

Introduction to University Mathematics

Preamble

The goal of this course is to introduce you to a range of mathematical ideas that are fundamental to studying degree-level mathematics. The course does not cover anything in great depth and is not foundational in the sense that we start from a list of axioms; that will happen in other courses. Rather, this course aims to provide a rapid introduction to various concepts, notation, and methods of logical reasoning, which you should find helpful as you begin studying mathematics at university.

These lecture notes are heavily based on material from previous notes authored by Richard Earl, Ian Hewitt, Alan Lauder, and Peter Neumann.

The colourful styling is enabled by the `thmtools` TeX package by Ulrich M. Schwarz, copying the set-up used by Evan Chen's Napkin project.

Please send corrections/queries to james.munro@maths.ox.ac.uk

James Munro, 2026

Synopsis

The natural numbers and their ordering. Induction as a method of proof, including a proof of the binomial theorem with non-negative integral coefficients.

Sets. Examples including $\mathbb{N}$, $\mathbb{Z}$, $\mathbb{Q}$, $\mathbb{R}$, $\mathbb{C}$, and intervals in $\mathbb{R}$. Inclusion, union, intersection, power set, ordered pairs and Cartesian product of sets. Relations. Definition of an equivalence relation. Examples.

Functions: composition, restriction; injective (one-to-one), surjective (onto) and invertible functions; images and preimages.

Writing mathematics. The language of mathematical reasoning; quantifiers: 'for all', 'there exists'. Formulation of mathematical statements with examples.

Proofs and refutations: standard techniques for constructing proofs; counter-examples. Example of proof by contradiction and more on proof by induction.

Problem-solving in mathematics: experimentation, conjecture, confirmation, followed by explaining the solution precisely.

Contents

0	The natural numbers and induction	3
0.0	The natural numbers	3
0.1	Mathematical induction	4
0.2	The binomial theorem	9
1	Sets	12
1.0	Definitions and examples	12
1.1	Algebra of sets	17
1.2	Truth tables	20
1.3	Cardinality	20
2	Relations	23
2.0	Definition and examples	23
2.1	Reflexivity, symmetry, anti-symmetry, and transitivity	24
2.2	Equivalence relations, equivalence classes, and partitions	25
3	Functions	28
3.0	Definitions and examples	28
3.1	Injectivity and surjectivity	31
3.2	Composition of functions and invertibility	33
4	Logic and Proof	37
4.0	Logical statements and notation	37
4.1	Implies	38
4.2	Formulation of mathematical statements	40
4.3	Examples of proof	44
4.4	Examples of problem-solving	47
4.5	General advice	50
A	Notes for two in-person lectures	52
A.0	Lecture 0	52
A.1	Lecture 1	53
B	The Greek alphabet	56
C	List of Definitions	57

0 The natural numbers and induction

0.0 The natural numbers

We start by discussing the natural numbers, which we define in the following way:

Definition 0.1 (Natural numbers)

A natural number is a member of the sequence $0, 1, 2, 3, \dots$, obtained by starting from 0 and adding 1 successively.

Notation

We write $\mathbb{N} = \{0, 1, 2, 3, \dots\}$ for the set of all natural numbers.

The curly bracket notation here indicates that the objects are grouped together as a *set*. A set is simply a collection of objects; we will discuss more about sets later in chapter 1.

Not everyone agrees that 0 should be included as a natural number. When counting (for example, when numbering sections of a document) most people tend to start at 1. Definitions are important in mathematics, and if you're reading material by two different authors then you should always check the definitions that they have each chosen to use.

Natural numbers can be added and multiplied. That is, if m and n are natural numbers, then we can construct $m + n$ and $m \times n$, which are also natural numbers. Addition and multiplication are examples of *binary operations*: they take a pair of elements from $\mathbb{N}$ and produce an element of $\mathbb{N}$. You've done so much adding and multiplying before in your life that you probably don't think that these need a definition. Nevertheless, we will attempt to define them below!

Two important natural numbers are 0 and 1, which are the additive and multiplicative identities, meaning they have the properties

$$n + 0 = n \quad \text{and} \quad n \times 1 = n \quad \text{for all } n \in \mathbb{N}.$$

The symbol $\in$ is shorthand to mean “is an element of”, and we read this as “for all n in the natural numbers”, or “for all natural numbers n ”.

Another important property of the natural numbers is that they have an *ordering*, so we can write things like $m \leq n$. We can carefully define this less-than-or-equal-to symbol:

Definition 0.2 (Less than or equal to)

Let m and n be natural numbers. We write $m \leq n$ to mean that there exists a natural number k such that $m + k = n$.

This is an example of a *relation*, which we'll discuss in chapter 2. This definition includes the case where $m = n$, because our definition of the natural numbers included 0, and because $m + 0 = m$.

Notice that our description of the natural numbers is a little unsatisfactory because it relies on the idea of “adding 1”, and we still haven't defined addition yet. For the purpose of this course we will rely on our basic intuition for these things, but it is possible to be more careful,

building on an axiomatic description of $\mathbb{N}$ (that is, laying out some clear *axioms* – statements that we assume to be true as a starting point – and then deducing all other properties from those). A popular way to do this is with the Peano axioms, one of which involves a “successor function”. Functions will be discussed in chapter 3.

0.1 Mathematical induction

We now move on to talk about induction. The following principle is sometimes quoted as a theorem, although it follows directly from our definition of the natural numbers. In fact it can be used as an axiom when defining $\mathbb{N}$ in a more rigorous manner.

Theorem 0.3 (Principle of Induction)

Let $P(n)$ be a family of statements indexed by the natural numbers. Suppose that

- (i) $P(0)$ is true, and
- (ii) for any n , if $P(n)$ is true then $P(n + 1)$ is also true.

Then $P(n)$ is true for all natural numbers n .

Induction is often visualised like toppling dominoes. The *inductive step* (ii) corresponds to placing each domino sufficiently close that it will be hit when the previous one falls over, and the *initial step*, (i) – often called the *base case* – corresponds to knocking over the first one.

To use induction to prove a family of statements, we simply have to demonstrate (i) and (ii). Before I show you an example, I would like to show you some notation for sums.

Notation

The expression

$$\sum_{k=a}^b f(k)$$

where f is a function means $f(a) + f(a + 1) + f(a + 2) + \cdots + f(b - 1) + f(b)$. We take the sum of the values of f over all integer k such that $a \leq k \leq b$.

Note that the sum is not itself a function and doesn’t involve k .

When written in-line, these sums are sometimes written like this; $\sum_{k=0}^n f(k)$.

By convention, empty sums like $\sum_{k=0}^{-3} f(k)$ are understood to be zero.

Notation

If we want the sum over something other than consecutive values of k , then we could write

$$\sum_{k \in S} f(k)$$

to indicate a sum over values in the set S , or we could even use words like

$$\sum_{p \text{ prime}} f(p),$$

but I'm getting ahead of myself.

We're ready for an example of proof by induction now.

Proposition 0.4

For any $n \in \mathbb{N}$,

$$\sum_{k=0}^n k = \frac{n(n+1)}{2}.$$

Proof. In this case $P(n)$ is the statement that the given equality holds for that particular n .

Clearly $P(0)$ holds because for $n = 0$ the sum on the left-hand side is 0 and the expression on the right-hand side is also 0.

Now suppose $P(n)$ holds. Then

$$\begin{aligned} \sum_{k=0}^{n+1} k &= \sum_{k=0}^n k + (n+1) \\ &= \frac{n(n+1)}{2} + (n+1) && \text{(by the inductive hypothesis)} \\ &= \frac{(n+1)(n+2)}{2}, \end{aligned}$$

which is exactly the statement $P(n+1)$. So by induction, $P(n)$ is true for all n . □

The small square on the right here is used to signify that we've reached the end of the proof. Historically the letters QED were used for this purpose, but that has largely gone out of fashion. Other symbols, including a filled square ■, are sometimes used.

The abbreviations LHS and RHS are commonly used to refer to the “left-hand side” and “right-hand side” of an equality. When using them in your mathematical writing, make sure it is clear which equality you are referring to. I've written a note at the side where I've used the inductive hypothesis (the assumption that $P(n)$ holds) to help communicate my argument. You should do the same; a well-written proof should always explain any steps that are not just routine algebra. You should always write in full sentences.

A straightforward extension of induction is if the family of statements holds for $n \geq N$, rather than necessarily $n \geq 0$.

Corollary 0.5

Let N be an integer and let $P(n)$ be a family of statements indexed by integers $n \geq N$. Suppose that

- (i) $P(N)$ is true, and
- (ii) for any $n \geq N$, if $P(n)$ is true then $P(n + 1)$ is also true.

Then $P(n)$ is true for all $n \geq N$.

Proof. This follows directly by applying the principle of induction to the statements $Q(n) = P(n + N)$ for $n \in \mathbb{N}$. $\square$

We use the word “corollary” to mean a result that is an extension of, or a consequence of, a theorem or proposition; a corollary is generally not such a major result as the theorem or proposition itself.

The words “theorem” and “proposition” are used somewhat interchangeably to mean a result that one has proved (unlike a “conjecture”, which is something that has not yet been proven). Theorem is typically used for more significant results, and theorems are often given a specific name.

We also use the word “lemma”, to mean a result that is going to be useful in proving a later theorem or proposition. Lemmas are typically not such exciting or major results in themselves.

Another variant on induction is when the inductive step relies on some earlier case(s) but not necessarily just the immediately previous case. This is sometimes called *strong induction*:

Theorem 0.6 (Strong Form of Induction)

Let $P(n)$ be a family of statements indexed by the natural numbers. Suppose that

- (i) $P(0)$ is true, and
- (ii) for any n , if $P(0), P(1), \dots, P(n)$ are true then $P(n + 1)$ is also true.

Then $P(n)$ is true for all natural numbers n .

Proof. Like above, we define a related family of statements $Q(n)$. To do this, let $Q(n)$ be the statement “ $P(k)$ holds for $k = 0, 1, \dots, n$ ”. Then the conditions for the strong form of induction are equivalent to (i) $Q(0)$ holds and (ii) for any n , if $Q(n)$ is true then $Q(n + 1)$ is also true. It follows by the (normal) principle of induction that $Q(n)$ holds for all n , and hence $P(n)$ holds for all n . $\square$

For an example to illustrate how the strong form of induction can be useful, I need to define a couple of terms that are hopefully familiar to you.

Definition 0.7 (Divides)

Given natural numbers m and n , we say that m divides n if there exists a natural number

k such that $m \times k = n$. We write $m|n$ if m divides n .

Definition 0.8 (Prime number)

A natural number $p > 1$ is prime if no natural numbers except 1 and p divide p .

Proposition 0.9

Every natural number greater than 1 may be expressed as a product of one or more prime numbers.

Proof. Let $P(n)$ be the statement that n may be expressed as a product of prime numbers.

Clearly $P(2)$ holds, since 2 is itself prime.

Let $n \geq 2$ be a natural number and suppose that $P(m)$ holds for all $m < n$. If n is prime then it is the product of one prime number, n . If n is not prime, then there must exist some $r, s > 1$ such that $n = rs$. By the inductive hypothesis, each of r and s can be written as a product of primes, and therefore $n = rs$ is also a product of primes. Thus, whether n is prime or not, $P(n)$ holds.

By strong induction, $P(n)$ is true for all natural numbers. That is, every natural number greater than 1 may be expressed as a product of one or more primes. $\square$

Related to induction is the idea of *recursion* as a method of definition. For example, if we suppose that we are happy with what it means to “add 1” (noting the discussion above), then we can recursively define more general addition on the natural numbers.

Definition 0.10 (Addition of natural numbers)

Define addition on $\mathbb{N}$ by the rules that for all $m \in \mathbb{N}$,

- (i) $m + 0 = m$, and
- (ii) for any $n \in \mathbb{N}$, $m + (n + 1) = (m + n) + 1$.

We can combine this with induction to prove some useful properties. For example,

Proposition 0.11 (Associativity)

Addition on $\mathbb{N}$ is associative. That is, for all $x, y, z \in \mathbb{N}$,

$$x + (y + z) = (x + y) + z.$$

Proof. We induct on z , so first suppose $z = 0$. Then, for any $x, y \in \mathbb{N}$,

$$\text{LHS} = x + (y + 0) = x + y = (x + y) + 0 = \text{RHS},$$

where we have twice used rule (i) from our definition of addition.

Now for the inductive step, suppose the proposition is true for $z = n$, and consider the case

$z = n + 1$. Then, for any $x, y \in \mathbb{N}$,

$$\begin{aligned}
 \text{LHS} &= x + (y + (n + 1)) \\
 &= x + ((y + n) + 1) && \text{(rule (ii) from the definition)} \\
 &= (x + (y + n)) + 1 && \text{(rule (ii) from the definition)} \\
 &= ((x + y) + n) + 1 && \text{(inductive hypothesis)} \\
 &= (x + y) + (n + 1) && \text{(rule (ii) from the definition)} \\
 &= \text{RHS}.
 \end{aligned}$$

So, by induction, the expression holds for any $z \in \mathbb{N}$. Thus addition is associative. $\square$

We can use a similar approach to define multiplication and factorial.

Definition 0.12 (Multiplication of natural numbers)

Define multiplication on $\mathbb{N}$ by the rules that for all $m \in \mathbb{N}$,

- (i) $m \times 0 = 0$, and
- (ii) for any $n \in \mathbb{N}$, $m \times (n + 1) = (m \times n) + m$.

Definition 0.13 (Factorial)

Define factorial $n!$ on $\mathbb{N}$ by the rules that

- (i) $0! = 1$, and
- (ii) for any $n \in \mathbb{N}$, $(n + 1)! = n! \times (n + 1)$.

Here is another important property of the natural numbers that we can prove using induction.

Theorem 0.14 (Well-ordering property of the natural numbers)

Every non-empty subset of $\mathbb{N}$ has a least element.

We have not yet defined a subset, but maybe you can guess what it means. S is a subset of $\mathbb{N}$ if every element of S is also an element of $\mathbb{N}$. Non-empty means it contains one or more elements.

Proof. We prove this by contradiction. Suppose, for a contradiction, that there is a non-empty subset S that does *not* have a least element. We define S^* to be the set of natural numbers that are not in S , and aim to show by induction that in fact $S^* = \mathbb{N}$.

Let $P(n)$ be the statement that S^* contains n .

For the initial step, note that 0 is not in S , since if it were, then S would have a least element (namely 0). So $0 \in S^*$ and therefore $P(0)$ holds.

Now suppose $P(0), \dots, P(n)$ hold. Then $n + 1$ cannot be in S , because if it were then it would be the least element of S (since by the inductive hypothesis all the smaller elements of $\mathbb{N}$ are not in S). Hence $n + 1 \in S^*$, and therefore $P(n + 1)$ holds.

The following result is an algebraic expression of the defining feature of Pascal's triangle, that each entry is the sum of the two entries most immediately above it.

Lemma 0.16 (Pascal's Triangle)

Let n and k be natural numbers with $1 \leq k \leq n$. Then

$$\binom{n}{k-1} + \binom{n}{k} = \binom{n+1}{k}.$$

Proof. This is simply a question of putting in the definitions and playing with the algebra. Putting the left hand side over a common denominator we obtain

$$\begin{aligned} \frac{n!}{(k-1)!(n-k+1)!} + \frac{n!}{k!(n-k)!} &= \frac{n!(k + (n-k+1))}{k!(n-k+1)!} \\ &= \frac{n! \times (n+1)}{k!(n-k+1)!} \\ &= \frac{(n+1)!}{k!(n+1-k)!}, \end{aligned}$$

which is equal to the right-hand side. □

This lemma can be used to show by induction that the binomial coefficients are integers rather than just rational numbers (a fact that is perhaps not immediately obvious from the definition).

We are now in a position to prove the binomial theorem.

Theorem 0.17 (Binomial Theorem)

Let x and y be real (or complex) numbers, and n be any natural number. Then

$$(x + y)^n = \sum_{k=0}^n \binom{n}{k} x^k y^{n-k}.$$

Proof. We use induction on n . First we check that the expression holds for $n = 0$. This is true, since the left hand side is 1 in that case, and the right hand side is also 1 (because $\binom{0}{0} = 1$ and any number raised to the power 0 is 1). Now assume the expression holds for n and consider the case for $n + 1$,

$$\begin{aligned} (x + y)^{n+1} &= (x + y)(x + y)^n \\ &= (x + y) \left(\sum_{k=0}^n \binom{n}{k} x^k y^{n-k} \right) \quad (\text{by the inductive hypothesis}). \end{aligned}$$

Continuing to expand the brackets gives

$$\sum_{k=0}^n \binom{n}{k} x^{k+1} y^{n-k} + \sum_{k=0}^n \binom{n}{k} x^k y^{n+1-k} = x^{n+1} + \sum_{k=0}^{n-1} \binom{n}{k} x^{k+1} y^{n-k} + \sum_{k=1}^n \binom{n}{k} x^k y^{n+1-k} + y^{n+1},$$

where we have taken out the last term from the first sum and the first term from the second sum. In the first sum we now make a change of indexing variable; we set $k = l - 1$, noting that as k ranges over $0, 1, \dots, n - 1$ then l ranges over $1, 2, \dots, n$. So the above equals

$$\begin{aligned}
 & x^{n+1} + \sum_{l=1}^n \binom{n}{l-1} x^l y^{n+1-l} + \sum_{k=1}^n \binom{n}{k} x^k y^{n+1-k} + y^{n+1} \\
 &= x^{n+1} + \sum_{k=1}^n \left\{ \binom{n}{k-1} + \binom{n}{k} \right\} x^k y^{n+1-k} + y^{n+1} && \text{[relabeling } l \text{ as } k] \\
 &= x^{n+1} + \sum_{k=1}^n \binom{n+1}{k} x^k y^{n+1-k} + y^{n+1} && \text{[using Lemma 0.16]} \\
 &= \sum_{k=0}^{n+1} \binom{n+1}{k} x^k y^{n+1-k},
 \end{aligned}$$

which shows that the expression holds for $n + 1$. Thus, by induction, the expression holds for all n . $\square$

1 Sets

1.0 Definitions and examples

We will rely on the intuitive idea that a set is a collection of objects. The objects are called the **elements**, or members, of the set. We might ask what exactly is meant by a “collection”, or by “objects”, and there are problems with taking too broad an interpretation of these terms. We *can* make things more precise, and you could take the Part B Set Theory course if you want to explore this more deeply.

Notation

For a set S , we write $x \in S$ to mean that x is an element of S , and we write $x \notin S$ to mean that x is not an element of S .

I usually read the expression “ $x \in S$ ” as “ x in S ”.

Definition 1.1 (Equal sets)

Two sets S and T are equal if and only if they contain the same elements, and in that case we can write $S = T$.

Definition 1.2 (Empty set)

The empty set is the set with no elements and is denoted by $\emptyset$.

The symbol $\emptyset$ is not to be confused with the Greek letter phi, written as ϕ or φ . This notation, which was inspired by the Danish or Norwegian symbol $\emptyset$, was introduced by André Weil (who is himself not to be confused with Andrew Wiles).

Definition 1.3 (Subset)

A set A is said to be a subset of a set S if every element of A is an element of S . We write $A \subseteq S$.

The symbol $\subseteq$ can be read as “is a subset of” or “is contained in”.

Definition 1.4 (Proper subset)

If $A \subseteq S$ and $A \neq S$, we call A a proper subset of S .

The symbol $\subset$ is commonly used to mean the same thing as $\subseteq$, and I’m afraid that it does not necessarily imply that the subset is proper (as you might otherwise imagine by analogy with $\leq$ and $<$ symbols).

You might see the symbol $A \not\subseteq S$ used to indicate that A is not a subset of S , and $A \subsetneq S$ used if A is a proper subset of S , with that notation trying to capture both parts of the definition; A is a subset of S , but A is not equal to S .

As we have already seen, curly brackets are used to denote sets.

Notation

The set with elements $a_1, a_2, \dots, a_n$ is written $\{a_1, a_2, \dots, a_n\}$.

If nothing is written at the end of a string of dots, such as $\{a_1, a_2, \dots\}$, it typically indicates an infinite list of elements, although this is not a hard and fast rule. It is helpful practice to write the final element, as done here, to indicate if such a list is finite.

The order of the elements inside the curly brackets is not important, and repeating the elements does not mean anything extra, so $\{3, 5\} = \{5, 3\}$, and $\{2, 3\} = \{2, 3, 3\}$.

We will often want to define a set in terms of some property $P(x)$ that the elements x satisfy. Typically the set is a subset of some larger set S , say.

Notation

We write $\{x \in S : P(x)\}$, or $\{x \in S \mid P(x)\}$, to mean the set of elements $x \in S$ that have property $P(x)$.

This is read as “the set of x in S such that P holds”.

Example 1.5

- (i) The set of even natural numbers is $\{n \in \mathbb{N} : n \text{ is divisible by } 2\}$.
- (ii) The set $\{n \in \mathbb{N} : n^2 < 0\}$ is equal to the empty set $\emptyset$, since no elements of $\mathbb{N}$ satisfy the given property.
- (iii) The set $\{n \in \mathbb{N} : 3|n, 5|n, 7|n\}$ is an example where multiple conditions are separated by commas, and all conditions must hold for a value of n to be included in the subset. This set is the multiples of 105 (prove it!), so I could have written $\{n \in \mathbb{N} : 105|n\}$.

Do not confuse a with $\{a\}$. The first is the element a and the second is the set containing the single element a . It’s easy to get confused if we start talking about sets of sets. For example, if $a = \emptyset$, the empty set, then a is a set with no elements, but $\{a\}$ is a set with one element (namely $\emptyset$).

Example 1.6

Check that you agree that each of the following sets has exactly two elements;

$$A = \{\mathbb{N}, 3\}, \quad B = \{\{1, 2\}, \{1, 3\}, \{1, 2\}\}, \quad C = \{\emptyset, \{\emptyset\}\}.$$

Note that it is correct to write $\mathbb{N} \in A$ and incorrect to write $\mathbb{N} \subseteq A$.

Note that $1 \notin B$.

Note that $\emptyset \in C$ and also $\emptyset \subseteq C$ (in fact, for any set S , $\emptyset \subseteq S$).

We have already seen the natural numbers, $\mathbb{N} = \{0, 1, 2, \dots\}$, as an example of a set. Some other important examples are:

Notation

The set of integers, $\mathbb{Z} = \{\dots, -2, -1, 0, 1, 2, \dots\}$, consisting of 0, the non-zero natural numbers (also known as the positive integers), and the negative integers $-1, -2, -3, -4, \dots$ made by writing a minus sign in front of each of the positive integers.

This is notation rather than a definition because, done properly, I would need to explain how to add and multiply these negative integers. Note that $\mathbb{N} \subset \mathbb{Z}$.

Notation

The set of rational numbers (or simply “rationals”), $\mathbb{Q}$, is the set comprising all fractions where the numerator m and denominator n are both integers. That is,

$$\mathbb{Q} = \left\{ \frac{m}{n} : m \in \mathbb{Z}, n \in \mathbb{Z}, n > 0 \right\}.$$

This is close to a definition, except I haven’t captured the idea that two fractions can be equivalent to each other like $\frac{1}{2}$ and $\frac{3}{6}$. Chapter 2 on equivalence relations will give us a way to describe this. Note that the rationals of the form $\frac{m}{1}$ are (loosely speaking) a copy of $\mathbb{Z}$, and so $\mathbb{Z} \subset \mathbb{Q}$.

Notation

The set of real numbers, $\mathbb{R}$, is the set containing numbers with a decimal expansion. These will be formally introduced in Analysis I.

You will notice that I’m avoiding the question of how to define these. The real numbers include the rational numbers (and so $\mathbb{Q} \subset \mathbb{R}$), but also *irrational* numbers such as $\sqrt{2}$ and π .

Notation

The set of complex numbers, $\mathbb{C} = \{a + bi : a, b \in \mathbb{R}\}$, with the multiplication rule $i^2 = -1$.

Note that, if we identify a number of the form $a + 0i \in \mathbb{C}$ with $a \in \mathbb{R}$, then we have $\mathbb{R} \subset \mathbb{C}$.

The symbols for each of these sets are written in ‘blackboard bold’ font. When writing them by hand there is no need to make them look quite so fancy; simply adding an extra line somewhere in the capital letter is sufficient.

Some frequently occurring subsets of the real numbers are **intervals**, which can be visualised as sections of the real line.

Definition 1.7 (Bounded interval, unbounded interval, open interval, closed interval)

Given real numbers a, b with $a \leq b$ we define bounded intervals

$$(a, b) = \{x \in \mathbb{R} : a < x < b\},$$

$$[a, b] = \{x \in \mathbb{R} : a \leq x \leq b\},$$

$$[a, b) = \{x \in \mathbb{R} : a \leq x < b\},$$

$$(a, b] = \{x \in \mathbb{R} : a < x \leq b\},$$

and unbounded intervals

$$\begin{aligned}(a, \infty) &= \{x \in \mathbb{R} : a < x\}, \\ [a, \infty) &= \{x \in \mathbb{R} : a \leq x\}, \\ (-\infty, a) &= \{x \in \mathbb{R} : x < a\}, \\ (-\infty, a] &= \{x \in \mathbb{R} : x \leq a\}.\end{aligned}$$

An interval of the form (a, b) or (a, ∞) or $(-\infty, a)$ is called an open interval.

An interval of the form $[a, b]$ or $[a, \infty)$ or $(-\infty, a]$ is called a closed interval.

Sadly, “closed” and “open” are not mutually exclusive. The interval $(a, b]$ is neither an open interval nor a closed interval.

If $a = b$, then $[a, b] = \{a\}$, while (a, b) and $[a, b)$ and $(a, b]$ are all $\emptyset$.

The notation above uses the ∞ symbol, which you might have seen before to represent the concept of “infinity”. Please note that $\infty \notin \mathbb{R}$ so, technically, speaking, it’s not included in our definition of $<$, and so we can’t write “ $x < \infty$ ”. However, many people do write “ $x < \infty$ ” to mean “ x is finite”.

The diagram below visualises three intervals as sections of the real line, using filled circles at the endpoints to indicate that the number there is included, and an empty circle to indicate that it is not.



We have seen some examples where the elements of the set are also sets. An important example of this is given by the following definition.

Definition 1.8 (Power set)

The power set of a set A , denoted $\mathcal{P}(A)$, is the set of all subsets of A .

Example 1.9

- (i) For $A = \{0, 1\}$, we have $\mathcal{P}(A) = \{\emptyset, \{0\}, \{1\}, \{0, 1\}\}$.
- (ii) For $A = \mathbb{N}$, $\mathcal{P}(A)$ has elements including the set of all even numbers, the set $\{0, 1\}$, the set of all primes p such that $p + 2$ is also prime, the empty set, and many more.

If we have two objects a and b we can combine them together to make a set $\{a, b\}$, but we could also combine them to make an **ordered pair** (a, b) . The distinction is that in an ordered pair the order matters. Two pairs (a_1, b_1) and (a_2, b_2) are equal if and only if $a_1 = a_2$ and $b_1 = b_2$. Unlike the set $\{a, a\}$, which we consider equal to the set $\{a\}$, for an ordered pair (a, a) is not considered equal to (a) . Indeed, (a) is not even considered to be an ordered pair.

You are probably familiar with such objects in the context of vectors or coordinates. If the elements are real numbers, for example, then (a, b) is a two-dimensional vector, and can be thought of as representing the x - y coordinates of a point on a plane (in that context it is clear

that $(1, 2)$ is quite different from $(2, 1)$; order matters!). Although we will often use ordered pairs in this way, the definition is much more general, since there is no reason why the elements a and b need come from the same sets or even be similar ‘types’ of object.

Definition 1.10 (Cartesian product)

Given sets A and B , the Cartesian product, denoted $A \times B$, is the set of all ordered pairs with the first element of the pair coming from A and the second from B . That is,

$$A \times B = \{(a, b) : a \in A \text{ and } b \in B\}.$$

Cartesian products are named after René Descartes.

Notation

If $A = B$, we might write $A \times A$ as A^2 .

Example 1.11

The case where $A = B = \mathbb{R}$ is a particularly important one that you will see in the Geometry course. In that case, the Cartesian product is $\mathbb{R}^2$ (usually read as “r two” rather than “r squared”), and represents the two-dimensional plane.

We can similarly create ordered triples (a, b, c) , quadruples (a, b, c, d) and so on. If there are n elements it is called an n -tuple.

Definition 1.12 (Cartesian product of n sets)

We define $A_1 \times A_2 \times \dots \times A_n$ to be the set of all ordered n -tuples $(a_1, a_2, \dots, a_n)$, where $a_i \in A_i$ for $1 \leq i \leq n$. If all the A_i are copies of A , we might write this Cartesian product as A^n .

You will have noticed that, regrettably, the same notation is used for open intervals (a, b) and ordered pairs (a, b) . The context usually helps!

Remark. The following example is known as Russell’s paradox (after the mathematician and philosopher Bertrand Russell, 1872–1970). It provides a warning as to the looseness of our definition of a set. Suppose

$$H = \{\text{sets } S : S \notin S\}.$$

That is, H is the collection of sets S that are not elements of themselves. All the sets we have come across seem to be in H (for example, $\mathbb{N}$ is in H since the elements of $\mathbb{N}$ are individual numbers and clearly none of them is the set $\mathbb{N}$ itself). The problem arises when we ask the question of whether or not H is itself in H ?

On the one hand, if $H \notin H$ then H meets the precise criterion for being in H and so $H \in H$, a contradiction. On the other hand, if $H \in H$ then by the property required for this to be the case, $H \notin H$, another contradiction. Thus we have a paradox: H seems to be neither in H nor not in H .

The modern resolution of Russell's Paradox is that we have taken too naïve an understanding of “collection”, and that Russell's “set” H is in fact not a set. It does not fit within axiomatic set theory (which relies on the so-called ZF axioms), and so the question of whether or not H is in H simply doesn't make sense.

1.1 Algebra of sets

Definition 1.13 (Union, intersection, complement, set difference)

Given subsets A and B of a set S , the union $A \cup B$ is the set consisting of those elements that are in A or B (or both), that is:

$$A \cup B = \{x \in S : x \in A \text{ or } x \in B\}.$$

The intersection $A \cap B$ is the set consisting of those elements that are in both A and B , that is:

$$A \cap B = \{x \in S : x \in A \text{ and } x \in B\}.$$

The complement of A , written A^c or sometimes A' , is the subset consisting of those elements that are not in A , that is:

$$A^c = \{x \in S : x \notin A\}.$$

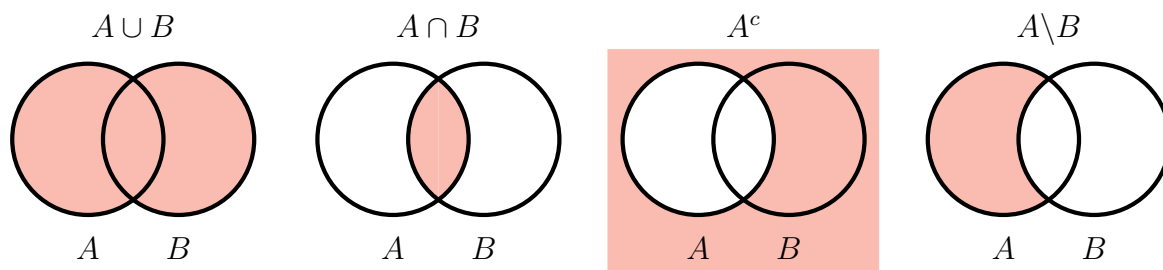
The set difference, or complement of B in A , written $A \setminus B$, is the subset of A consisting of those elements that are not in B , that is:

$$A \setminus B = \{x \in A : x \notin B\}.$$

Note that $A \setminus B = A \cap B^c$.

Note that $A \setminus B$ uses a backslash, not to be confused with a forward slash. The notation A/B is used in Group Theory with a different meaning entirely.

A useful way of visualising these is using a Venn diagram. The following is an example, where A and B are the regions inside the two circles, respectively:



Definition 1.14 (Disjoint sets)

Two sets A and B are said to be disjoint if $A \cap B = \emptyset$, that is the two subsets have no element in common.

More generally, we can take unions and intersections of arbitrary numbers of sets, even infinitely many. If we have a family of subsets $\{A_i : i \in I\}$, where I is called an **indexing set**, we write

$$\bigcap_{i \in I} A_i = \{x \in S : x \in A_i \text{ for all } i \in I\},$$

and

$$\bigcup_{i \in I} A_i = \{x \in S : x \in A_i \text{ for at least one } i \in I\}.$$

Often I might be a subset of $\mathbb{N}$, in which case we write things like

$$\bigcap_{i=1}^n A_i, \quad \bigcup_{i=1}^{\infty} A_i,$$

but it could be an even “larger” set such as $\mathbb{R}$.

Example 1.15

Let S be the set of all students at Oxford, $A \subseteq S$ be the set of students studying mathematics, and $B \subseteq S$ be the set of students at your college. Then $A \cap B$ is the set of students studying mathematics at your college, $A \cup B$ is the set of students either at your college or studying mathematics (this is probably the set where many of your friends will come from), B^c is the set of all students at other colleges, $A \setminus B$ is the set of students studying mathematics at other colleges.

Example 1.16

The set of irrational numbers is $\mathbb{R} \setminus \mathbb{Q}$.

Example 1.17

Some authors define the set of natural numbers to be what I would write as $\mathbb{N} \setminus \{0\}$. Note that I have to wrap the 0 in curly braces, because 0 is not a set and $A \setminus B$ is only defined when B is a set.

The following result may well seem obvious, but it provides quite an important recipe for how to show that two sets are the same, as we will see below.

Proposition 1.18 (Double Inclusion)

Let A and B be two subsets of a set S . Then $A = B$ if and only if $A \subseteq B$ and $B \subseteq A$.

Proof. If $A = B$, then every element in A is an element in B , so certainly $A \subseteq B$, and similarly $B \subseteq A$.

Conversely, suppose $A \subseteq B$, and $B \subseteq A$. Then for every element $x \in S$, if $x \in A$ then $A \subseteq B$ implies that $x \in B$, and if $x \notin A$ then $B \subseteq A$ means $x \notin B$.

So $x \in A$ if and only if $x \in B$, and therefore $A = B$. □

Here is an example of how we can make use of double inclusion to show that two sets are equal:

Proposition 1.19 (Distributive Laws)

Let A, B, C be subsets of a set S . Then

$$(i) \quad A \cup (B \cap C) = (A \cup B) \cap (A \cup C)$$

$$(ii) \quad A \cap (B \cup C) = (A \cap B) \cup (A \cap C)$$

Proof. We first prove (i).

(LHS $\subseteq$ RHS): Suppose x is in the LHS of (i), that is $x \in A \cup (B \cap C)$. This means that $x \in A$ or $x \in B \cap C$. Thus either $x \in A$ or x is in both B and C . If $x \in A$ then $x \in A \cup B$ and $x \in A \cup C$, and therefore x is in the RHS. If x is in both B and C then similarly x is in both $A \cup B$ and $A \cup C$. Thus every element of the LHS is in the RHS, which means we have shown $A \cup (B \cap C) \subseteq (A \cup B) \cap (A \cup C)$.

(RHS $\subseteq$ LHS): Conversely suppose that $x \in (A \cup B) \cap (A \cup C)$. Then x is in both $A \cup B$ and $A \cup C$. Thus either $x \in A$ or, if $x \notin A$, then $x \in B$ and $x \in C$. Thus $x \in A \cup (B \cap C)$. Hence $(A \cup B) \cap (A \cup C) \subseteq A \cup (B \cap C)$.

By double inclusion, $(A \cup B) \cap (A \cup C) = A \cup (B \cap C)$.

The proof of (ii) follows similarly and is left as an exercise. □

In the proof above there were two separate things to show (that is, LHS $\subseteq$ RHS, and RHS $\subseteq$ LHS, which combine to give the required result). When laying out a proof like this it is helpful to separate the two things out clearly, both to aid your own understanding and that of the reader. Here we have done that by starting a new paragraph, while the smaller individual steps of the argument were written as sentences. When hand-writing a proof, some people tend to put each step of the logic on a new line, in which case leaving a larger gap or using different indentation may help to distinguish the separate sections. In the proof above I have labelled the sections with bracketed comments like “(LHS $\subseteq$ RHS)”.

Here is an example of how set complements work.

Proposition 1.20 (De Morgan’s Laws)

Let A and B be subsets of a set S . Then

$$(A \cup B)^c = A^c \cap B^c \quad \text{and} \quad (A \cap B)^c = A^c \cup B^c.$$

Proof. For the first one, suppose $x \in (A \cup B)^c$. Then x is not in either A or B . Thus $x \in A^c$ and $x \in B^c$, and therefore $x \in A^c \cap B^c$.

Conversely, suppose $x \in A^c \cap B^c$. Then $x \notin A$ and $x \notin B$, so x is in neither A nor B , and therefore $x \in (A \cup B)^c$.

By double inclusion, the first result holds.

The second result follows similarly and is again left as an exercise. □

De Morgan’s laws extend naturally to any number of sets, so if $\{A_i : i \in I\}$ is a family of

subsets of S , then

$$\left(\bigcap_{i \in I} A_i\right)^c = \bigcup_{i \in I} A_i^c \quad \text{and} \quad \left(\bigcup_{i \in I} A_i\right)^c = \bigcap_{i \in I} A_i^c.$$

1.2 Truth tables

Another way of proving set-theoretical identities is via **truth tables**. These provide a systematic way of cataloguing all the different cases for whether or not a given element is in each set. Here is an example:

A	B	$A \cap B$	$A \cup B$	$A \setminus B$
F	F	F	F	F
F	T	F	T	F
T	F	F	T	T
T	T	T	T	F

The different cases are listed in the rows, one row for each, and various sets are placed along the columns. We put a T or an F in each column to indicate whether it is true or false that the given element is in that set for each case. There will be different numbers of cases to consider depending on the number of sets involved; in this example with two sets, there are four (the first two columns effectively define these four different cases, and the entries in the other columns then follow from those).

Truth tables like this can be used to catalogue the cases of ‘true’ or ‘false’ for a series of logical statements. These truth tables for sets are a particular instance, where the statements are of the form ‘ $x \in A$ ’, ‘ $x \in B$ ’, ‘ $x \in A \cap B$ ’, and so on.

Here is an alternative proof of De Morgan’s laws using a truth table:

Proof of Proposition 1.20. We list the four combinations of cases for whether or not $x \in A$ and $x \in B$:

A	B	$A \cap B$	$(A \cap B)^c$	$A \cup B$	$(A \cup B)^c$	A^c	B^c	$A^c \cup B^c$	$A^c \cap B^c$
F	F	F	T	F	T	T	T	T	T
F	T	F	T	T	F	T	F	T	F
T	F	F	T	T	F	F	T	T	F
T	T	T	F	T	F	F	F	F	F

Comparing columns, the fact that $(A \cap B)^c$ and $A^c \cup B^c$ are the same in every case shows that those two sets are the same. Similarly, the fact that $(A \cup B)^c$ and $A^c \cap B^c$ are the same in every case shows that those two sets are the same. $\square$

Truth tables quickly become unwieldy for longer statements, and I prefer to write out cases separately, for example “if $x \in A$ then... and if $x \notin A$ then...”.

1.3 Cardinality

Informally, the cardinality of a set S , denoted $|S|$, is a measure of its ‘size’. For finite sets, there is little ambiguity about this – it is simply the number of distinct elements in the set (we give a formal definition below, which will also clarify what it means to say that a set is finite).

But for infinite sets, things are more interesting. For example, one might be tempted to think that the set of even natural numbers is in some sense ‘smaller’ than the set of natural numbers – there might reasonably seem to be ‘fewer’ of them, since we have left the odd numbers out. But if we simply divide every element in that set by 2, then we see that for each even natural number we have a natural number, and vice versa (multiplying by 2) for each natural number we have an even natural number. By this logic it seems that these two sets ought to be the same size. Indeed, these two sets do have the same cardinality, $\aleph_0$ (pronounced ‘aleph-null’). This is the smallest infinite cardinal. The concept of cardinals was invented and investigated widely by Georg Cantor, 1845–1918. Perhaps surprisingly, the rational numbers $\mathbb{Q}$ also have the same cardinality $\aleph_0$. But the cardinality of the real numbers $\mathbb{R}$ is larger, as will be discussed more in the Analysis I course.

We will be able to give a nicer definition of cardinality later, once we have discussed bijections, but the following provides a recursive definition of the cardinality for a finite set.

Definition 1.21 (Finite set (recursive definition))

A set S is finite if either

- (i) it is the empty set $\emptyset$, or
- (ii) there exists $s \in S$ such that $S \setminus \{s\}$ is finite.

Any set that is not finite is said to be **infinite**.

Definition 1.22 (Cardinality of a finite set (recursive definition))

The cardinality of a finite set S , written $|S|$, is a natural number defined as follows;

- (i) if $S = \emptyset$ then $|S| = 0$.
- (ii) if S is not empty, then by definition of finiteness there exists $s \in S$ such that $S \setminus \{s\}$ is finite. If $|S \setminus \{s\}| = n$ then we define $|S| = n + 1$.

It is not hard to see that this means that if $S = \{x_1, x_2, \dots, x_n\}$, and $x_i \neq x_j$ whenever $i \neq j$, then $|S| = n$. Conversely, if $|S| = n$ then S is a set with n elements.

Proposition 1.23

Let A and B be finite sets. Then $|A \cup B| = |A| + |B| - |A \cap B|$.

The proof is left as an exercise (see problem sheet).

Proposition 1.24 (Subsets of a finite set)

If a set A is finite with $|A| = n$, then its power set is finite and has $|\mathcal{P}(A)| = 2^n$.

Proof. We use induction. For the initial step, note that if $|A| = 0$ then $A = \emptyset$ has no elements, so there is a single subset, $\emptyset$, and therefore $|\mathcal{P}(A)| = 1 = 2^0$.

Now suppose that $n \geq 0$ and that $|\mathcal{P}(S)| = 2^n$ for any set S with $|S| = n$. Let A be

any set with $|A| = n + 1$. By definition, this means that there is an element a and a set $A' = A \setminus \{a\}$ with $|A'| = n$. Any subset of A must either contain the element a or not, so we can partition $\mathcal{P}(A) = \mathcal{P}(A') \cup \{S \cup \{a\} : S \in \mathcal{P}(A')\}$. These two sets are disjoint, and each of them has cardinality $|\mathcal{P}(A')| = 2^n$ by the inductive hypothesis. Hence, by Proposition 1.23, $|\mathcal{P}(A)| = 2^n + 2^n = 2^{n+1}$.

Thus, by induction, the result holds for all n . $\square$

An alternative, perhaps easier, way to see why the size of the power set should be $2^{|A|}$ for a finite set A , is to consider the process of creating a subset. We can do this systematically by going through each of the $|A|$ elements in A and making the yes/no decision whether to put it in the subset. Since there are $|A|$ such choices, that yields $2^{|A|}$ different combinations of elements and therefore $2^{|A|}$ different subsets.

2 Relations

2.0 Definition and examples

In mathematics a **relation** (sometimes *binary relation*) is something like $\leq$ or $\subseteq$ or $=$ that tells us that there is a relationship between two objects.

Formally, we define it as the set of ordered pairs (a, b) for which the relation holds.

Definition 2.1 (Relation)

A relation R on a set S is a subset of $S \times S$.

Notation

If $(a, b) \in R$, we write aRb .

It's a little bit awkward that we use the letter R for both the name of the relation and the symbol to indicate that a is related to b , and indeed in many cases there is some other conventional notation that we use to indicate that a is related to b .

It may seem odd to think of $\leq$ as a subset like this, but this definition turns out to be a convenient way to describe such a relation as a mathematical object. We could alternatively think of a relation as a binary operation that takes two input elements and returns a “True” or “False”, depending on whether the two elements are related according to that relation. But this is effectively the same thing as our definition above, since the subset R is simply those elements of $S \times S$ that return “True” under that operation.

Example 2.2

- (i) The “less than or equal to” relation on the set of real numbers is $\{(x, y) \in \mathbb{R}^2 : x \leq y\}$. We write $x \leq y$ if (x, y) is in this set.
- (ii) The “divides” relation on $\mathbb{N}$ is $\{(m, n) \in \mathbb{N}^2 : m \text{ divides } n\}$. We write $m|n$ if (m, n) is in this set.
- (iii) For a set S , the “subset” relation on $\mathcal{P}(S)$ is $\{(A, B) \in \mathcal{P}(S)^2 : A \subseteq B\}$. We write $A \subseteq B$ if (A, B) is in this set.

I would also like you to be aware of some trivial examples, not because these are particularly interesting themselves, but because they can be helpful as we start to think about properties that relations do or do not have.

Example 2.3

For any set S , we can define the following relations.

- (i) *Everything's related*. There is a relation R defined by “for all x and for all y , xRy ”.
- (ii) *Nothing's related*. There is a relation R defined by “for all x , there does not exist y such that xRy ”.

- (iii) *Everything's related only to itself.* There is a relation R defined by “for all x , xRy if and only if $y = x$.”

Note that if $S = \emptyset$, then these all describe the same relation, giving us perhaps the ultimate trivial relation.

Finally, let's have a real-world example.

Example 2.4

Consider the set C of all chess players. Then “ aRb if and only if a has ever beaten b at chess” is a relation on C .

Note that for this relation, aRb and bRa mean different things, and those things are not mutually exclusive. Also, there are pairs of elements a and b such that neither aRb nor bRa . For example, Magnus Carlsen has never beaten me at chess, and I haven't beaten him (yet).

2.1 Reflexivity, symmetry, anti-symmetry, and transitivity

Definition 2.5 (Reflexive, symmetric, anti-symmetric, transitive)

Let S be a set, R a relation on S and $x, y, z \in S$. We say that

- (i) R is reflexive if xRx for all x in S ,
- (ii) R is symmetric if whenever xRy then yRx ,
- (iii) R is anti-symmetric if whenever xRy and yRx then $x = y$,
- (iv) R is transitive if whenever xRy and yRz then xRz .

Example 2.6

For the three trivial examples, decide whether they have the properties we just defined (does it matter what the set S is?)

Convince yourself that the relation in the chess example above satisfies none of these properties.

The relation $\leq$ on $\mathbb{R}$ is reflexive, anti-symmetric, and transitive.

The relation $<$ on $\mathbb{R}$ is not reflexive or symmetric, but it is anti-symmetric and transitive.

The relation $\neq$ on $\mathbb{R}$ is not reflexive, anti-symmetric or transitive, but it is symmetric.

The relation $=$ on $\mathbb{R}$ is reflexive, symmetric, and transitive.

Here is another important example:

Example 2.7

Let $n \geq 2$ be an integer, and define R on $\mathbb{Z}$ by saying aRb if and only if $a - b$ is a multiple of n . Then R is reflexive, symmetric and transitive.

Proof. Reflexivity: For any $a \in \mathbb{Z}$ we have aRa as 0 is a multiple of n .

Symmetry: If aRb then $a - b = kn$ for some integer k . So $b - a = -kn$, and hence bRa .

Transitivity: If aRb and bRc then $a - b = kn$ and $b - c = ln$ for integers k, l . So then $a - c = (a - b) + (b - c) = (k + l)n$, and hence aRc . $\square$

2.2 Equivalence relations, equivalence classes, and partitions

Example 2.7 provides an example of a particularly important type of relation, an **equivalence relation**. An equivalence relation provides a way of saying two objects are, in some particular sense, ‘the same’:

Definition 2.8 (Equivalence relation)

A relation R on a set S is an equivalence relation if it is reflexive, symmetric and transitive.

Notation

If R is an equivalence relation, we denote it by $\sim$ (various other symbols, including $\equiv$, are sometimes used).

Example 2.9

The following are all examples of equivalence relations:

- (i) $S = \mathbb{C}$, with $z \sim w$ if and only if $|z| = |w|$.
- (ii) S is the set of polygons in $\mathbb{R}^2$, and $\sim$ is congruence.
- (iii) S is the set of differentiable functions on $\mathbb{R}$, and $f \sim g$ if and only if $f'(x) = g'(x)$ for all x .
- (iv) The relation on $\mathbb{Z}$ defined in Example 2.7; in this case $\sim$ represents congruence modulo n . It is the basis for modular arithmetic, and you will often see $a \sim b$ expressed as $a \equiv b \pmod{n}$ in this case.

An equivalence relation provides a way of grouping together elements that can be viewed as being the same.

Definition 2.10 (Equivalence class)

Given an equivalence relation $\sim$ on a set S , and given $x \in S$, the equivalence class of x , denoted $\bar{x}$ (or sometimes $[x]$), is the subset

$$\bar{x} = \{y \in S : y \sim x\}.$$

For example, with the equivalence relation defined in Example 2.7 (congruence modulo n) the equivalence class of 1 is the set $\bar{1} = \{\dots, -n+1, 1, n+1, 2n+1, \dots\}$; that is, all the integers that are congruent to 1 modulo n . Note that 1 and $\bar{1}$ are different.

Definition 2.11 (Quotient set)

Given an equivalence relation $\sim$ on a set S , the set of equivalence classes is called the quotient set and is denoted $S/\sim$. We have

$$S/\sim = \{\bar{x} : x \in S\}$$

or more explicitly

$$S/\sim = \{\{y \in S : y \sim x\} : x \in S\}.$$

Note that $S/\sim$ is a set of sets.

Grouping the elements of a set into equivalence classes like this provides a partition of the set, which we define as follows.

Definition 2.12 (Partition)

A partition of a set S is a collection of subsets $\{A_i \subseteq S : i \in I\}$, where I is an indexing set, with the property that

- (i) $A_i \neq \emptyset$ for all $i \in I$ (that is, all the subsets are non-empty),
- (ii) $\bigcup_{i \in I} A_i = S$ (that is, every member of S lies in one of the subsets),
- (iii) $A_i \cap A_j = \emptyset$ for every $i \neq j$ (that is, the subsets are disjoint).

The subsets are called the parts of the partition.

For example, $\{n \in \mathbb{N} : n \text{ is divisible by } 2\}, \{n \in \mathbb{N} : n + 1 \text{ is divisible by } 2\}$ forms a partition of the natural numbers, into evens and odds.

The fact that the equivalence classes for any equivalence relation form a partition is something that will be proved in the Groups and Group Action course later in the year (though it is not particularly difficult and you may like to have a go at doing so).

Conversely, we can use any given partition to define an equivalence relation, by saying that $x \sim y$ if and only if x and y are elements of the same part of the partition (you may like to check that indeed this definition satisfies the conditions to be an equivalence relation: reflexivity, symmetry and transitivity). Thus, there is a natural correspondence between equivalence relations and partitions of a set.

Example 2.13

Suppose S is the set of students at Oxford. This set can be partitioned according to colleges; that is, the partitioning subsets are the sets of students at each college, which form a partition since (i) there are no colleges with no students (I'm ignoring All Souls College for the purpose of this example), (ii) every student is a member of a college, and

(iii) you can't be at more than one college (I'm ignoring some exceptions). Then the equivalence relation $\sim$ induced by this partition says that $x \sim y$ if and only if x and y are at the same college.

This is a good example of the fact that being “equivalent” does not mean the elements are actually the same! Every one of you is wonderfully unique... but there might be some purpose for which it is convenient to view all the students in a college as essentially the same, and the equivalence class provides a way of representing that mathematically.

3 Functions

3.0 Definitions and examples

Definition 3.1 (Function)

Let X and Y be sets. A function f from X to Y is a subset of $X \times Y$ with the property that for each $x \in X$, there is exactly one value of $y \in Y$ such that (x, y) is in the subset.

For example, with $X = \{1, 2, 3\}$ and $Y = \{a, b, c\}$, the set $\{(1, a), (2, b), (3, a)\}$ is a function from X to Y .

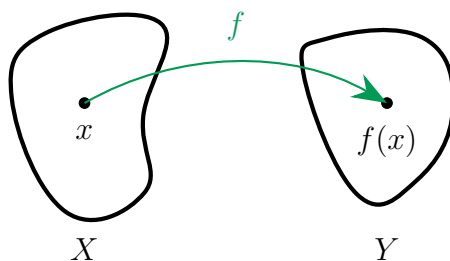
Of course, this is probably not how you think of functions in practice (and indeed, some people would say that what I've defined above is in fact the **graph** of a function). The property in the definition means that we can think of the function as a **map** or **mapping** from X to Y such that for each $x \in X$ there is an assignment of a particular element $f(x) \in Y$.

With this way of thinking, a function takes an input element from the set X , and maps it to an output in the set Y . The input is often referred to as the *argument* of the function (the word 'argument' has many different meanings!).

Notation

We write $f : X \rightarrow Y$ to mean that f is a function from X to Y .

The arrow here highlights the idea that the function takes values *from* X and *to* Y .



Definition 3.2 (Domain, codomain)

For a function $f : X \rightarrow Y$, the set X is called the domain of f , and the set Y is called the codomain of f .

Definition 3.3 (Equal functions)

Two functions $f : X \rightarrow Y$ and $g : X \rightarrow Y$ are equal (in which case we can write $f = g$), if and only if $f(x) = g(x)$ for every $x \in X$.

Note that we require functions to have the same domain as each other, and also the same codomain as each other, in order to be equal.

Example 3.4

We can define a function $f: \mathbb{R} \rightarrow \mathbb{Z}$ that takes any real number x and returns the least integer that is greater than or equal to x . (This is called the **ceiling** function, and is denoted by $\lceil x \rceil$. You may be familiar with it as “rounding up”; there is a related **floor** function, $\lfloor x \rfloor$ that ‘rounds down’).

Then $f(3.5) = 4$ and $f(-4) = -4$ but, for example, $f(1.01 + 2.5i)$ is not defined.

Notation

We might define a particular function by explaining what each element x maps to, using the notation $x \mapsto f(x)$

Example 3.5

The function that adds one to a natural number can be defined as $f: \mathbb{N} \rightarrow \mathbb{N}$ with $n \mapsto n + 1$.

Beware; the function f is not the same thing as the value $f(x)$. This matters when we start to talk about the properties of the function f itself, which might be very different from properties of values of the function $f(x)$.

The definition of a function requires that a unique element of the codomain is assigned for every element of the domain, so when we define a function we need to take account of this. For example if we want to define a function $f: \mathbb{R} \rightarrow \mathbb{R}$, the assignment $x \mapsto 1/x$ is not sufficient, since it fails at $x = 0$. Similarly, the recipe “ $f(x) = y$ where $y^2 = x$ ” fails for two reasons: one is that $f(x)$ is undefined for $x < 0$, and the other is that for $x > 0$ it does not return a unique value – does $f(4)$ equal 2 or -2 ? In such cases, we say the “function” is **ill-defined** (it doesn’t satisfy the definition of a function). We are interested in the opposite; functions that are **well-defined**. When we start to define more complicated functions, it may not be so immediately obvious whether this is the case, so some effort is often required to demonstrate that a given recipe produces a well-defined function.

Another example of what might seem to be a function, but which is not well-defined (and therefore not actually a function), is if we try to let $f: \mathbb{Q} \rightarrow \mathbb{Z}$ be given by $f(\frac{m}{n}) = n$. The problem here is that there is not a unique way of expressing an element of $\mathbb{Q}$ as $\frac{m}{n}$, so this assignment gives multiple different values for the same argument (for example $f(\frac{2}{3}) = 3$ and $f(\frac{4}{6}) = 6$, but for the function to be well-defined we need these to be the same).

As the remarks above highlight, the definition of a function needs to make clear the domain and codomain, not just the “formula”. The following are all different functions.

Example 3.6

$$\begin{array}{lll}
 f_1: \mathbb{R} \rightarrow \mathbb{R} & \text{given by} & f_1(x) = x^2. \\
 f_2: \mathbb{R} \rightarrow [0, \infty) & \text{given by} & f_2(x) = x^2. \\
 f_3: [0, \infty) \rightarrow \mathbb{R} & \text{given by} & f_3(x) = x^2. \\
 f_4: [0, \infty) \rightarrow [0, \infty) & \text{given by} & f_4(x) = x^2.
 \end{array}$$

Definition 3.7 (Image, range, pre-image)

Given a function $f: X \rightarrow Y$, the image or range of f is

$$f(X) = \{f(x) : x \in X\} \subseteq Y.$$

More generally, given $A \subseteq X$, the image of A (under f) is

$$f(A) = \{f(x) : x \in A\} \subseteq Y.$$

Given $B \subseteq Y$, the pre-image of B (under f) is

$$f^{-1}(B) = \{x : f(x) \in B\} \subseteq X.$$

Beware that some authors say “range” for what I would call the codomain of a function.

Example 3.8

For the function f_1 defined above in Example 3.6, the image is $[0, \infty)$, and $f_1([0, 1]) = [0, 1]$, $f_1^{-1}([0, 1]) = [-1, 1]$, and $f_1^{-1}((-\infty, 0]) = \{0\}$.

For f_4 , the image is $[0, \infty)$, and $f_4([0, 1]) = [0, 1]$, $f_4^{-1}([0, 1]) = [0, 1]$, and $f_4^{-1}((-\infty, 0])$ is not defined, since $(-\infty, 0]$ is not a subset of the codomain in this case.

Beware the confusing notation: for $x \in X$, $f(x)$ is a single *element* of Y , but for $A \subseteq X$, $f(A)$ is a *set* (a subset of Y). Formally, the function f induces a map of power sets $f_*: \mathcal{P}(X) \rightarrow \mathcal{P}(Y)$ defined by $f_*(A) = \{f(x) : x \in A\}$ for all $A \subseteq X$. But outside of Category Theory, this distinction is not usually made, and we’ll write f for both functions.

Beware also; the notation $f^{-1}(B)$ should be read as “the pre-image of B ” and not as “ f -inverse of B ”. The pre-image is defined even if no inverse function exists (in which case f^{-1} on its own has no meaning; we discuss invertibility of a function below). Also, we’ll see that even when the inverse function exists, it’s a function on Y , whereas this notation describes a function on the power set of Y . The notation $f^{-1}(B)$ for the pre-image of a set is widely used like this, I’m afraid, so we’ll have to work with it.

Proposition 3.9

Let $f: X \rightarrow Y$ be a function.

- (a) For any $A \subseteq X$ we have $A \subseteq f^{-1}(f(A))$, but do not have equality in general.
- (b) For any $C \subseteq Y$ we have $f(f^{-1}(C)) \subseteq C$, but do not have equality in general.

Proof. (a) By definition $A \subseteq f^{-1}(f(A))$ as the elements of A map into $f(A)$. However other elements may also map into $f(A)$. If we consider the map $f(x) = x^2$ from $\mathbb{R}$ to $\mathbb{R}$ then we see for $A = \{1\}$ that

$$f^{-1}(f(A)) = f^{-1}(\{1\}) = \{-1, 1\} \neq \{1\} = A.$$

(b) We immediately have $f(f^{-1}(C)) \subseteq C$ as f maps the elements of $f^{-1}(C)$ into C by definition. However $f(f^{-1}(C))$ need not equal C . For example with the same f as in (a) and $C = \{-1\}$ we see that $f^{-1}(C) = \emptyset$, so

$$f(f^{-1}(C)) = f(\emptyset) = \emptyset \neq C.$$

□

If a function is defined on some larger domain than we care about, it may be helpful to restrict the domain:

Definition 3.10 (Restriction of a function)

Given a function $f: X \rightarrow Y$ and a subset $A \subseteq X$, the restriction of f to A is the map $f|_A: A \rightarrow Y$ defined by $f|_A(x) = f(x)$ for all $x \in A$.

The restriction is almost the same function as the original f – just the domain has changed.

Another rather trivial but nevertheless important function is the identity map:

Definition 3.11 (Identity function)

Given a set X , the identity $\text{id}_X: X \rightarrow X$ is defined by $\text{id}_X(x) = x$ for all $x \in X$. Other notation is sometimes used, such as 1_X , and if the domain is unambiguous, we might just write id .

3.1 Injectivity and surjectivity

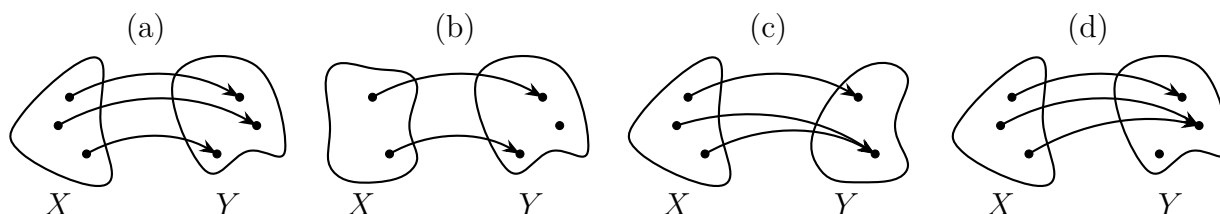
Definition 3.12 (Injective, surjective, bijective, bijection)

Let $f: X \rightarrow Y$ be a function.

- (i) We say that f is injective, or one-to-one, if whenever $f(x_1) = f(x_2)$ then $x_1 = x_2$. In words, each element of the image has a unique pre-image.
- (ii) We say that f is surjective, or onto, if for every $y \in Y$ there exists $x \in X$ such that $f(x) = y$. In words, the image of the domain is the codomain.
- (iii) We say that f is bijective if it is both injective and surjective. In words, each element of the codomain has a unique pre-image.

A bijective function is called a bijection (and similarly, we say injection and surjection for injective and surjective functions respectively).

The figure below shows schematic representations of functions $f: X \rightarrow Y$ that are (a) both injective and surjective, (b) injective but not surjective, (c) surjective but not injective, (d) neither injective nor surjective. Note that X and Y are not the same between pictures.



Example 3.13

The ceiling function from $\mathbb{R}$ to $\mathbb{Z}$ defined in Example 3.4 is surjective, because for any integer $n \in \mathbb{Z}$, we have $\lceil n \rceil = n$. But it is not injective, because there are other elements in $\mathbb{R}$ (infinitely many of them in fact) that will return the same output; for example $\lceil n - \frac{1}{2} \rceil = n$.

Example 3.14

Returning to the functions defined in Example 3.6, we see that the function f_1 is not injective (because, for example, $f(-1) = f(1) = 1$), and it is not surjective (because there are no $x \in \mathbb{R}$ that give $x^2 = y$ for $y < 0$).

The function f_2 is also not injective but *is* surjective (because $y < 0$ is not in the codomain this time, and for every $y \in [0, \infty)$ there *is* some $x \in \mathbb{R}$ with $x^2 = y$).

The function f_3 is again not surjective, but this time it is injective (because negative values are now excluded from the domain).

The function f_4 is both injective and surjective (and is therefore a bijection).

Having introduced injective and surjective functions, we can give an alternative and more intuitive definition of the cardinality of finite sets:

Definition 3.15 (Finite set, cardinality of a set (with bijections))

The empty set $\emptyset$ is finite and has cardinality $|\emptyset| = 0$. A non-empty set S is said to be finite and have cardinality $|S| = n \in \mathbb{N}$ if and only if there exists a bijection from S to the set $\{1, 2, \dots, n\}$.

The bijection provides a way of counting the elements of S . You might like to convince yourself that this definition is equivalent to the inductive definition given earlier.

Note that for finite sets X and Y , a function $f: X \rightarrow Y$ can only be injective if $|Y| \geq |X|$, since for any injective function the number of elements in the image $f(X)$, is equal to the number of elements in the domain, and $f(X) \subseteq Y$. In other words, the codomain of an injective function

cannot be smaller than the domain. This is sometimes referred to as the **pigeonhole principle** (so called from the observation that if n letters are placed in m pigeonholes and $n > m$, then at least one hole must contain more than one letter; the non-injective function in that case is the assignment of pigeonholes to letters).

Similarly, a function between finite sets $f: X \rightarrow Y$ can only be surjective if $|Y| \leq |X|$. Hence if f is bijective, then $|X| = |Y|$; that is, the domain and codomain of a bijection have equal cardinality.

For infinite sets, it is a version of this new ‘bijective’ definition that leads to the conclusion that $\mathbb{N}$ and $\mathbb{Z}$ and $\mathbb{Q}$ have the same cardinality (there are bijections between these sets), but $\mathbb{R}$ does not have the same cardinality (there is no bijection between $\mathbb{N}$ and $\mathbb{R}$).

3.2 Composition of functions and invertibility

Definition 3.16 (Function composition)

Given two functions $f: X \rightarrow Y$ and $g: Y \rightarrow Z$, the composition $g \circ f: X \rightarrow Z$ is defined by

$$(g \circ f)(x) = g(f(x)) \quad \text{for all } x \in X.$$

If $Z = X$ then we can similarly define $f \circ g: Y \rightarrow Y$, but in general $f \circ g \neq g \circ f$ (indeed, if $X \neq Y$, these functions have different domain and codomain, but even in the case $X = Y$ the functions will generally not be the same). For example, if $f(x) = x^2$ and $g(x) = e^x$ are both maps from $\mathbb{R}$ to $\mathbb{R}$, then

$$(f \circ g)(x) = e^{2x} \neq e^{x^2} = (g \circ f)(x).$$

This shows that composition of functions is not *commutative*. However, composition is *associative*, as the following result shows:

Proposition 3.17 (Function composition is associative)

Let $f: X \rightarrow Y$, $g: Y \rightarrow Z$, $h: Z \rightarrow W$ be three functions. Then

$$h \circ (g \circ f) = (h \circ g) \circ f.$$

Proof. Let $x \in X$. Then, by the definition of composition, we have

$$(h \circ (g \circ f))(x) = h((g \circ f)(x)) = h(g(f(x))) = (h \circ g)(f(x)) = ((h \circ g) \circ f)(x).$$

□

The following proposition addresses the extent to which composition of functions preserves injectivity and surjectivity:

Proposition 3.18

Let $f: X \rightarrow Y$ and $g: Y \rightarrow Z$ be functions.

- (i) If f and g are injective then so is $g \circ f$. Conversely, if $g \circ f$ is injective, then f is injective, but g need not be.
- (ii) If f and g are surjective then so is $g \circ f$. Conversely, if $g \circ f$ is surjective, then g is surjective, but f need not be.

Proof and commentary. We prove (i), and leave the proof of (ii) as an exercise.

It is helpful to clarify for each part of the proposition what exactly we are told (the hypotheses), and what exactly we need to show. For the first part of (i), we can take it that f and g are injective, and need to show that $g \circ f$ is injective. From the definition of injectivity, that means we need to show that for any $x_1, x_2 \in X$, if $(g \circ f)(x_1) = (g \circ f)(x_2)$ then $x_1 = x_2$. So let's suppose $x_1, x_2 \in X$ and $(g \circ f)(x_1) = (g \circ f)(x_2)$, and aim to show $x_1 = x_2$, making use of what we know. From the injectivity of g we know that if $g(f(x_1)) = g(f(x_2))$ then $f(x_1) = f(x_2)$, so this must be the case here. Then from the injectivity of f we know that this means $x_1 = x_2$. So we have indeed shown what is needed for $g \circ f$ to be injective.

For the second part of (i), we are told that $g \circ f$ is injective, and we need to show that f is injective; that is, we need to show that if $f(x_1) = f(x_2)$ then $x_1 = x_2$. So let's suppose that $f(x_1) = f(x_2)$ and aim to show this, making use of what we know about $g \circ f$ this time. Applying g to both sides gives $g(f(x_1)) = g(f(x_2))$, and then we see that the injectivity of $g \circ f$ immediately tells us that $x_1 = x_2$. So we have shown that f is injective.

An alternative approach here could have been to use contradiction. If we start with the supposition that f is *not* injective, then it means there exist some $x_1, x_2 \in X$ for which $x_1 \neq x_2$ but $f(x_1) = f(x_2)$. Then $g(f(x_1)) = g(f(x_2))$, so in fact this means there exist some $x_1, x_2 \in X$ for which $x_1 \neq x_2$ but $(g \circ f)(x_1) = (g \circ f)(x_2)$. But that would contradict the definition of $g \circ f$ being injective, so our supposition that f was not injective was incorrect, and we have therefore shown that f *is* injective.

To show that g need not be injective, we should give a counterexample. A bit of thought may lead to the observation that g could have a larger domain than the image of f . An extreme example is to take $X = Z = \{0\}$ and $Y = \mathbb{R}$, and have f and g defined by $f(0) = 0$ and $g(y) = 0$ for all $y \in \mathbb{R}$. Then $(g \circ f): X \rightarrow Z$ is injective (it simply maps 0 to 0). But clearly g is not injective. $\square$

I have written much more than is needed in the proof above because I am spelling out the thought process as well as the logic. Having worked out what to do, we could streamline it:

Proof. For the first part of (i), suppose $(g \circ f)(x_1) = (g \circ f)(x_2)$ for some $x_1, x_2 \in X$. From the injectivity of g we know that $g(f(x_1)) = g(f(x_2))$ implies $f(x_1) = f(x_2)$, and then from the injectivity of f we know that this implies $x_1 = x_2$. So $g \circ f$ is injective.

For the second part of (i), suppose $f(x_1) = f(x_2)$ for some $x_1, x_2 \in X$. Then applying g gives $g(f(x_1)) = g(f(x_2))$, and by the injectivity of $g \circ f$ this means $x_1 = x_2$. So f is injective. To see that g need not be injective, a counterexample is $X = Z = \{0\}$, $Y = \mathbb{R}$, with $f(0) = 0$ and $g(y) = 0$ for all $y \in \mathbb{R}$. $\square$

Recalling that id_X is the identity map on a set X , we are now in a position to define invertibility:

Definition 3.19 (Invertible function, inverse function)

A function $f: X \rightarrow Y$ is invertible if there exists a function $g: Y \rightarrow X$ such that $g \circ f = \text{id}_X$ and $f \circ g = \text{id}_Y$. The function g is the inverse of f , and we write $g = f^{-1}$.

Note that directly from the definition, if f is invertible then f^{-1} is also invertible, and $(f^{-1})^{-1} = f$.

An immediate concern we might have is whether there could be multiple such functions g , in which case the inverse f^{-1} would not be well-defined. This is resolved by the following result:

Proposition 3.20

If $f: X \rightarrow Y$ is invertible then its inverse is unique.

Proof. Let g_1 and g_2 be two functions for which $g_i \circ f = \text{id}_X$ and $f \circ g_i = \text{id}_Y$. Using the fact that composition is associative, and the definition of the identity maps, we can write

$$g_1 = g_1 \circ \text{id}_Y = g_1 \circ (f \circ g_2) = (g_1 \circ f) \circ g_2 = \text{id}_X \circ g_2 = g_2.$$

□

As a result, we can describe a rather non-trivial example of a function.

Example 3.21

Let X and Y be sets, and let F be the set of all invertible functions with domain X and codomain Y . Then there is a function $T: F \rightarrow F$ such that $T(f) = f^{-1}$, and this function T is its own inverse, with $T \circ T = \text{id}_F$.

The following result shows how to invert the composition of invertible functions.

Proposition 3.22

Let $f: X \rightarrow Y$ and $g: Y \rightarrow Z$ be functions between sets X, Y, Z . If f and g are invertible, then $g \circ f$ is invertible, and $(g \circ f)^{-1} = f^{-1} \circ g^{-1}$.

Proof. Making repeated use of the fact that function composition is associative, and the definition of the inverses f^{-1} and g^{-1} , we note that

$$\begin{aligned} (f^{-1} \circ g^{-1}) \circ (g \circ f) &= ((f^{-1} \circ g^{-1}) \circ g) \circ f \\ &= (f^{-1} \circ (g^{-1} \circ g)) \circ f \\ &= (f^{-1} \circ \text{id}_Y) \circ f \\ &= f^{-1} \circ f \\ &= \text{id}_X, \end{aligned}$$

and similarly,

$$\begin{aligned}
 (g \circ f) \circ (f^{-1} \circ g^{-1}) &= g \circ (f \circ (f^{-1} \circ g^{-1})) \\
 &= g \circ ((f \circ f^{-1}) \circ g^{-1}) \\
 &= g \circ (\text{id}_Y \circ g^{-1}) \\
 &= g \circ g^{-1} \\
 &= \text{id}_Z,
 \end{aligned}$$

which shows that $f^{-1} \circ g^{-1}$ satisfies the properties required to be the inverse of $g \circ f$. $\square$

The following result provides an important and useful criterion for invertibility:

Theorem 3.23

A function $f: X \rightarrow Y$ is invertible if and only if it is bijective.

Proof. First suppose f is invertible, so that it has an inverse $f^{-1}: Y \rightarrow X$. To show f is injective, suppose that for some $x_1, x_2 \in X$ we have $f(x_1) = f(x_2)$. Then applying f^{-1} to both sides and noting that by definition $f^{-1} \circ f = \text{id}_X$, we see that $x_1 = f^{-1}(f(x_1)) = f^{-1}(f(x_2)) = x_2$. So f is injective. To show that f is surjective, let $y \in Y$, and note that $f^{-1}(y) \in X$ has the property that $f(f^{-1}(y)) = y$. So f is surjective. Therefore f is bijective.

Conversely, suppose that f is bijective. We aim to show that there is a well-defined $g: Y \rightarrow X$ such that $g \circ f = \text{id}_X$ and $f \circ g = \text{id}_Y$. Since f is surjective, we know that for any $y \in Y$, there is an $x \in X$ such that $f(x) = y$. Furthermore, since f is injective, we know that this x is unique. So for each $y \in Y$ there is a unique $x \in X$ such that $f(x) = y$. This recipe provides a well-defined function $g(y) = x$, for which we have $g(f(x)) = x$ for any $x \in X$ and $f(g(y)) = y$ for any $y \in Y$. So g satisfies the property required to be an inverse of f and therefore f is invertible. $\square$

It is also possible to define left-inverse and right-inverse functions as functions that partially satisfy the definition of the inverse.

Definition 3.24 (Left-invertible, right-invertible)

A function $f: X \rightarrow Y$ is left-invertible if there exists a function $g: Y \rightarrow X$ such that $g \circ f = \text{id}_X$, and is right-invertible if there exists a function $h: Y \rightarrow X$ such that $f \circ h = \text{id}_Y$.

As may be somewhat apparent from the previous proof, being left- and right- invertible is equivalent to being injective and surjective, respectively. We leave this as an exercise to show.

Proposition 3.25

- (i) A function $f: X \rightarrow Y$ is left-invertible if and only if it is injective.
- (ii) A function $f: X \rightarrow Y$ is right-invertible if and only if it is surjective.

The proof is left as an exercise.

4 Logic and Proof

4.0 Logical statements and notation

We have already seen that we must deal with a lot of precise statements and assertions, ranging from very simple ones like “ $n = 2$ ” to slightly more involved ones like

$$\text{for all } x \in \mathbb{R}, x^2 \geq 0,$$

or

$$\text{there exist } x, y, z \in \mathbb{N} \setminus \{0\} \text{ such that } x^{1729} + y^{1729} = z^{1729}.$$

The statements can be true or false (perhaps you recognise the final example just given, which is famously false, by Fermat’s last Theorem). Often during a mathematical argument we might not (yet) know if a given statement is true or false. Nevertheless we can work with it; following through the logical consequences of a statement may eventually lead us to the conclusion that it is false, for example.

To discuss such logic it is helpful to denote the statements by symbols; we have already seen examples of this. If P is the statement ‘ $x \geq 2$ ’ and Q is ‘ $x^2 \geq 4$ ’ (in the context of $\mathbb{R}$), we can then say things like ‘ P implies Q ’ (which is, itself, another logical statement – one that we know to be true in this case).

We can combine logical statements using connecting words like ‘and’, and ‘or’, and we can negate a statement P by writing ‘not P ’.

For the avoidance of doubt, we define these with truth tables.

Definition 4.1 (And, Or, Not)

If P and Q are statements, then the statements ‘ P and Q ’ and ‘ P or Q ’ are true or false as per the following truth table.

P	Q	P and Q	P or Q
F	F	F	F
F	T	F	T
T	F	F	T
T	T	T	T

If P is a statement, then the statement ‘not P ’ is true or false as per the following truth table.

P	not P
F	T
T	F

In regular English, the word ‘or’ is often interpreted as an *exclusive* or; that is, it may carry an implicit meaning of ‘one but not the other’ (as in ‘you can have a piece of cake or an ice cream’). This is not the case in mathematical usage, where ‘ P or Q ’ should be interpreted to mean that P holds or Q holds or both do.

Symbols are sometimes used for these.

Notation

$P \wedge Q$ means ‘ P and Q ’. $P \vee Q$ means ‘ P or Q ’. $\neg P$ means ‘not P ’.

Example 4.2

If P is ‘ $n = 2$ ’ and Q is ‘ n is even’ (in the context of $\mathbb{N}$), then ‘ $P \vee Q$ ’ is equivalent to ‘ n is even’, and ‘ $P \wedge Q$ ’ means ‘ $n = 2$ ’, and $\neg P$ is the same as ‘ $n \neq 2$ ’.

There is a direct analogy between the symbols $\vee$, $\wedge$ and $\neg$, and the symbols $\cup$, $\cap$, and c for set union, intersection and complement. To make this analogy clear, suppose $A, B \subseteq S$ are sets and let P and Q be the statements ‘ $x \in A$ ’ and ‘ $x \in B$ ’. Then $x \in A \cup B$ is clearly equivalent to $P \vee Q$, $x \in A \cap B$ is equivalent to $P \wedge Q$, and $x \in A^c$ is equivalent to $\neg P$.

These logical symbols therefore obey the same distributive laws as for sets (Proposition 1.19) and also De Morgan’s laws (Proposition 1.20), which in this context are

Proposition 4.3 (De Morgan’s laws for logic)

- (i) $\neg(P \vee Q)$ if and only if $(\neg P) \wedge (\neg Q)$
- (ii) $\neg(P \wedge Q)$ if and only if $(\neg P) \vee (\neg Q)$.

These rules for how to negate ‘ P or Q ’ or ‘ P and Q ’ are hopefully quite intuitive (if we don’t have P or Q holding then that means we don’t have P holding *and* we don’t have Q holding). But when it comes to more complicated statements it is easy to confuse ourselves, so having these clear rules to fall back on may be useful.

4.1 Implies

Definition 4.4 (Implies)

The statement ‘ P implies Q ’ means $Q \vee (\neg P)$ and is written $P \Rightarrow Q$. With a truth table we can write

P	Q	$P \Rightarrow Q$
F	F	T
F	T	T
T	F	F
T	T	T

In regular English usage, ‘implies’ or ‘if ... then ...’ tends to be understood to indicate a degree of *causation*; that X has something to do with Y . In mathematical usage this does not need to be the case. So if we say ‘ $P \Rightarrow Q$ ’ or ‘If P then Q ’, we simply mean that whenever P is true, then Q is also true.

So ‘If Paris is the capital of France then the Thames flows through London’ is a true statement, despite the fact that there is obviously no connection between these two facts.

Similarly, ‘If Oxford is on Mars then Cambridge is on Venus’ is also a true statement. A statement like this, where the ‘ P ’ is never true, is said to be *vacuously true*. A similarly useless

statement might begin ‘for all $x \in \emptyset, \dots$ ’.

Note that, unlike $P \wedge Q$ and $P \vee Q$, there is an asymmetry here; ‘ P implies Q ’ is not the same statement as ‘ Q implies P ’.

Example 4.5

If P is ‘ $n = 2$ ’ and Q is ‘ n is even’, then $P \Rightarrow Q$, but $Q \not\Rightarrow P$.

Statements of the form $P \Rightarrow Q$ are very common, and there are lots of ways to say the same thing. The following are equivalent.

- (i) if P then Q ;
- (ii) P implies Q ;
- (iii) $P \Rightarrow Q$;
- (iv) P only if Q ;
- (v) P is a sufficient condition for Q ;
- (vi) Q is a necessary condition for P ;
- (vii) if Q does not hold then P does not hold;
- (viii) not Q implies not P ;
- (ix) $(\neg Q) \Rightarrow (\neg P)$.

The last three of these are known as the contrapositive.

Definition 4.6 (Contrapositive, converse)

- (i) For a statement of the form ‘ $P \Rightarrow Q$ ’, the contrapositive is $\neg Q \Rightarrow \neg P$. It is **always** true when the original statement is true.
- (ii) For a statement of the form $P \Rightarrow Q$, the converse is $Q \Rightarrow P$. It is **sometimes but not always** true when the original statement is true.
- (iii) For a statement of the form $P \Rightarrow Q$, the negation is $\neg(P \Rightarrow Q)$, which is **never** true when the original statement is true.

Using De Morgan’s laws for logic (Proposition 4.3), the negation of $P \Rightarrow Q$ can also be written as $P \wedge (\neg Q)$, that is ‘ P is true and Q is false’.

In order to prove a statement of this form, we typically start by assuming that P holds and try to deduce through some logical steps that Q holds too. Alternatively, we can start by assuming that Q does not hold and show that P does not hold (that is, we prove the contrapositive).

If using $\Rightarrow$, note that the symbol stands for both the ‘if’ and the ‘then’. We shouldn’t use it to stand for just the ‘then’ as, for example, in ‘if $x = -1 \Rightarrow x^2 = 1$ ’. This would mean ‘if

$x = -1$ implies that $x^2 = 1$ and would need to be followed by ‘then ...’ (to match the ‘if’).

As noted above, even if we know that $P \Rightarrow Q$, it might or might not be the case that $Q \Rightarrow P$.

Notation

We write $P \Leftrightarrow Q$ to mean ‘ $P \Rightarrow Q$ and $Q \Rightarrow P$ ’.

You can use the definition of $P \Rightarrow Q$ and De Morgan’s laws for logic (Proposition 4.3) to deduce that $P \Leftrightarrow Q$ means $(P \wedge Q) \vee (\neg P \wedge \neg Q)$.

We can read this as ‘ P if and only if Q ’, or ‘ P is necessary and sufficient for Q ’, or ‘ P is equivalent to Q ’.

The use of the $\Leftrightarrow$ symbol in a proof needs some caution, as we’ll discuss later (it has a logical implication both forwards and backwards, whereas our thought process tends to work in one direction most of the time!). The letters ‘iff’ are also commonly used to stand for ‘if and only if’.

These statements are usually best thought of separately as ‘if’ and ‘only if’ statements. So to prove ‘ P if and only if Q ’, we should first prove ‘if P then Q ’, and then separately prove ‘if Q then P ’ (or vice versa). Sometimes we may find that essentially the same argument used for the first direction also works in reverse, but sometimes quite a different method of argument may be required. One thing to be wary of is that having *assumed* P to deduce Q , and then having changed to assuming Q with a view to *deducing* P , it is all too easy to keep making use of P (or parts of P), forgetting that that is no longer assumed. It is a good idea to make very clear, both to yourself and in your written proof, which direction you are doing.

4.2 Formulation of mathematical statements

We’ve seen the phrases ‘for all’ and ‘there exists’ many times already in these lectures. It’s time to introduce some notation for them.

Notation

The symbol $\forall$ denotes ‘for all’ or ‘for every’, and can simply be used as shorthand for those words.

The symbol $\exists$ denotes ‘there exists’, and can similarly replace those words.

The symbols $\forall$ and $\exists$ are known as **quantifiers**.

Typically a phrase like ‘there exists $x \in S$ ’ is followed by a statement $P(x)$ about what specific property x has, and it is quite common for $\exists$ to stand for the following ‘such that’ as well as the ‘there exists’. For example, ‘ $\exists x \in S \quad P(x)$ ’ can be read as ‘there exists x in S such that $P(x)$ holds’. Personally I prefer to include the letters ‘s.t.’ (or a colon ‘:’) to stand in place of the ‘such that’, so I write ‘ $\exists x \in S$ s.t. $P(x)$ ’.

Notation

The symbol $\exists!$ means ‘there exists unique’, implying that there is one, and only one, element with the given property.

It is helpful to practice ‘translating’ such statements into English sentences, and vice versa.

Example 4.7

Using these symbols we can write things like

$$\forall x \in \mathbb{R} \quad x^2 \geq 0 \quad \text{but} \quad \exists x \in \mathbb{C} \text{ s.t. } x^2 < 0,$$

which you should read as ‘for all x in the real numbers, x^2 is greater than or equal to zero, but there exists x in the complex numbers such that x^2 is less than zero’.

Example 4.8

The definitions of what it means for $f: X \rightarrow Y$ to be injective and surjective can be written as (respectively)

$$\forall x_1, x_2 \in X, \quad f(x_1) = f(x_2) \Rightarrow x_1 = x_2,$$

and

$$\forall y \in Y, \quad \exists x \in X \text{ s.t. } f(x) = y.$$

Example 4.9

If $\mathbb{P}$ is the set of prime numbers, then

$$\forall p \in \mathbb{P}, \quad p > 2 \Rightarrow \exists n \in \mathbb{N} \text{ s.t. } p = 2n + 1,$$

is a way of stating ‘every prime number greater than 2 is odd’, and

$$\forall x \in \mathbb{R}, \quad (x < 0 \quad \text{or} \quad \exists y \in \mathbb{R} \text{ s.t. } y^2 = x),$$

is a way of stating ‘every non-negative real number has a real square root’.

The quantifiers should include a specification of the set over which they range ($\mathbb{P}$ or $\mathbb{R}$ in these examples). However, there are situations when this is so obvious from the context that it becomes cumbersome to keep writing this. In particular, you’ll often see ‘ $\forall \varepsilon > 0$ ’, in which it is understood that ε is a real number.

To prove a statement of the form ‘ $\forall x \in X \quad P(x)$ ’, it is always a good idea to start the proof with ‘Let $x \in X$.’ or ‘Suppose $x \in X$ is given.’. This ‘addresses’ the quantifier with an arbitrary x , which should then be treated as fixed for the rest of the proof.

It is important to note that the order of quantifiers matters. Returning to an earlier example, if S is the set of students at Oxford, and C is the set of colleges, we can say

$$\forall s \in S, \quad \exists c \in C \text{ s.t. } s \in c,$$

which says that for every student there is a college of which they are a member; this is true. If we changed the order of the quantifiers and wrote

$$\exists c \in C \text{ s.t. } \forall s \in S, \quad s \in c,$$

this says that there is one particular college of which every student at Oxford is a member; that is a completely different statement, and it is not true.

Importantly, we must read from left to right, and as new elements or statements are introduced they are allowed to depend on previously introduced elements but can't depend on things that are yet to be mentioned. So in the first of the statements above, the $c \in C$ that exists according to the second quantifier can (and does) depend on the specific $s \in S$ from the first quantifier. In the second statement, the specific c identified by the first quantifier must work for all s in the second one.

Proposition 4.10 (Negation)

- (i) The negation of a statement of the form ' $\forall x \in X, P(x)$ ' is ' $\exists x \in X$ s.t. $\neg P(x)$ '.
- (ii) The negation of a statement of the form ' $\exists x \in X$ s.t. $P(x)$ ' is ' $\forall x \in X, \neg P(x)$ '.

We can therefore negate a complicated statement like ' $\forall s \in S, \exists c \in C$ s.t. $s \in c$ ' by working from the outside in. Each line in the following is equivalent.

$$\begin{aligned} & \neg (\forall s \in S, \exists c \in C \text{ s.t. } s \in c) \\ & \exists s \in S \text{ s.t. } \neg (\exists c \in C \text{ s.t. } s \in c) \\ & \exists s \in S \text{ s.t. } \forall c \in C, \neg (s \in c) \\ & \exists s \in S \text{ s.t. } \forall c \in C, s \notin c \end{aligned}$$

If we try to express these in words, we might have something roughly like

It's not true that every student has a college.
 There's a student for whom it's not true that they have a college.
 There's a student, and for every college, it's not the case that they're a member.
 There's a student and for every college, they're not a member.

But one reason why logical notation is so helpful is that English itself is sometimes ambiguous. For example, the statement

‘For all natural numbers x , $x < y$ for some natural number y ’

could mean ‘for all $x \in \mathbb{N}$, there exists $y \in \mathbb{N}$ such that $x < y$ ’, or ‘there exists $y \in \mathbb{N}$ such that for all $x \in \mathbb{N}$, $x < y$ ’. The English wording could justifiably be interpreted either way, but clearly only the first of these is true. In symbolic notation, we should write

$$\forall x \in \mathbb{N}, \exists y \in \mathbb{N} \text{ s.t. } x < y,$$

and it is unambiguous that y is allowed to depend on x .

To avoid confusion, it is a good idea to keep to the convention that the quantifiers come first, before any statement to which they relate. However, many authors (including this one!) don't stick rigidly to this if there's a last ‘for all’. For example, if $f : \mathbb{R} \rightarrow \mathbb{R}$ is a bounded function, you may see something like

$$\exists M \in \mathbb{R} \text{ s.t. } |f(x)| < M \quad \forall x \in \mathbb{R}.$$

Example 4.11

As an example of how the notation introduced in this section can make a proof more concise, we'll probe the double inclusion principle again.

Proof of Proposition 1.18. We argue, via a sequence of equivalent statements, that $A = B$ is the same as $(A \subseteq B \text{ and } B \subseteq A)$:

$$\begin{aligned} A = B &\Leftrightarrow \forall x \in S \quad (x \in A \Leftrightarrow x \in B) \\ &\Leftrightarrow \forall x \in S \quad (x \in A \Rightarrow x \in B \quad \text{and} \quad x \in B \Rightarrow x \in A) \\ &\Leftrightarrow \forall x \in S \quad (x \in A \Rightarrow x \in B) \quad \text{and} \quad \forall x \in S \quad (x \in B \Rightarrow x \in A) \\ &\Leftrightarrow A \subseteq B \quad \text{and} \quad B \subseteq A. \end{aligned}$$

□

Notice that bracketing is sometimes necessary in order to make clear what the statements are to which the quantifiers refer. Intelligent use of spacing on the page can also be helpful, especially when writing by hand.

To use $\Leftrightarrow$ like this we need to make sure that the logic works both ways, both forwards and backwards. When constructing such a proof, it can be helpful to first work forwards, with each statement implying the next one, and then separately check whether the argument works backwards (i.e. first write it with the arrows as $\Rightarrow$, before checking if each one can be converted to $\Leftrightarrow$).

It is to some extent a matter of personal taste how much to use symbolic notation rather than writing things out in words. Regardless of how much you use them in your own writing, it is important to understand and be fluent in interpreting these symbols in other people's writing.

In order to prove or use a theorem it is important to correctly understand its logical form. Most theorems are ultimately of the form 'if P then Q ', although the P and the Q may themselves be quite complicated statements that are combinations of other statements. In this context, the ' P ' is the **hypothesis** and the ' Q ' is the **conclusion**.

Example 4.12

Consider the statement

If n is a non-zero natural number, then n has a unique prime factorisation.

Here, the hypothesis is ' n is a non-zero natural number', and the conclusion is ' n has a unique prime factorisation'.

In this example, the theorem is explicitly stated in the form 'if P then Q ', so understanding the hypothesis and conclusions is very easy. Sometimes a theorem may be stated in a way that makes this less obvious.

Example 4.13

Consider the statement

Every prime number greater than 2 is odd.

Faced with such a statement, it may be helpful to think through carefully what the hypothesis and conclusion are, and to re-state it in a way that makes this more transparent. Thus, another way of saying the same thing is:

Let p be a prime number greater than 2. Then p is odd.

The first sentence is the hypothesis, and the second is the conclusion.

Example 4.14

As a more involved example, here is a rather poor statement of the intermediate value theorem (IVT), which you will come across in Analysis.

Whenever f is a continuous function on $\mathbb{R}$, a, b are real numbers such that $a < b$, $f(a) < 0$ and $f(b) > 0$, $f(c) = 0$ for some $c \in (a, b)$.

Although all the correct ingredients of the theorem are here, it is not at all clear how to split the statement into hypothesis and conclusion. A better version is:

Let $f: \mathbb{R} \rightarrow \mathbb{R}$ be a continuous function, and suppose $a, b \in \mathbb{R}$ are such that $a < b$, $f(a) < 0$, and $f(b) > 0$. Then there exists a real number $c \in (a, b)$ such that $f(c) = 0$.

Splitting the theorem into shorter sentences has helped to clearly separate the hypothesis (which is a combination of various sub-statements in this case) from the conclusion. As a general rule, using words like ‘Let’ and ‘Suppose’ is a good way of ‘setting up’ the hypotheses, and ‘Then’ is a good way of signalling that what follows is the conclusion.

The intermediate value theorem is stating the intuitively obvious result that if the graph of a continuous function is below the axis somewhere and above the axis somewhere else, then it must *cross* the axis somewhere in between.

4.3 Examples of proof

We have already seen various methods for proving and refuting mathematical statements. We now make some comments on the general classes of methods. To illustrate the discussion, we’ll consider the following simple result about the arithmetic and geometric means of non-negative real numbers:

Theorem 4.15 (AM-GM Inequality)

Let x, y be non-negative real numbers. Then

$$\frac{x + y}{2} \geq \sqrt{xy},$$

with equality if and only if $x = y$.

The left hand side of this inequality is called the arithmetic mean (AM) and the right hand side is called the geometric mean (GM).

Direct proof

To prove a statement of the form ‘if P then Q ’ directly, we make use of P to arrive at Q through a sequence of logical reasoning. It may be that we can start from P and work directly to Q , or it may be that we make use of P along the way (as in the example below). A direct proof of Theorem 4.15 could look like this:

Proof of Theorem 4.15. Since $x, y \in \mathbb{R}$, we know that $(x - y)^2 \geq 0$, with equality if and only if $x = y$. Expanding the brackets, and then adding $4xy$ yields

$$\begin{aligned}x^2 - 2xy + y^2 &\geq 0, \\x^2 + 2xy + y^2 &\geq 4xy, \\ \frac{1}{4}(x + y)^2 &\geq xy.\end{aligned}$$

Taking the square root (noting that $x, y \geq 0$) gives the required result. $\square$

In this case, the algebraic steps are straightforward, but the starting point and the ‘adding $4xy$ ’ are perhaps not entirely obvious. You will see many direct proofs like this, and you must understand that the arguments are not necessarily presented in the order that the author thought of them. To illustrate this, let’s see another proof where the ideas are presented in a different order.

Proof by contradiction

To prove a statement ‘if P then Q ’ by contradiction, we suppose that Q is not true and show through some logical reasoning (making use of the hypotheses P) that this leads to a contradiction or inconsistency. We may arrive at something that contradicts the hypotheses P , or something that contradicts the initial supposition that Q is not true, or we may arrive at something that we know to be universally false.

A proof by contradiction of Theorem 4.15 could look like:

Proof of Theorem 4.15. Suppose for contradiction that $\exists x, y \in \mathbb{R}$ with $x, y \geq 0$ and $\sqrt{xy} > \frac{1}{2}(x + y)$.

Then squaring both sides (noting that both sides are positive) and rearranging, yields

$$\begin{aligned} 4xy &> 4x^2 + 2xy + y^2, \\ 0 &> x^2 - 2xy + y^2, \\ 0 &> (x - y)^2, \end{aligned}$$

which is a contradiction, since the square of a real number is non-negative. Hence $\sqrt{xy} \leq \frac{1}{2}(x + y)$. The same steps with $>$ replaced by $=$ show that equality holds if and only if $x = y$. $\square$

One of the useful aspects of proving things by contradiction is that by negating the statement Q , we immediately give ourselves something extra to work with. So if we cannot see a way to make any progress by starting directly from P , then supposing a contradiction may be a good thing to try instead. In the proof above, for example, we simply started with the negation of the result we were aiming for and manipulated the statement to find the contradiction.

Proof by induction

Induction is useful for proving results that can be indexed by the natural numbers. So it would not be useful as a method to prove Theorem 4.15, but it could be used to prove the generalisation:

Theorem 4.16 (AM-GM Inequality)

Let $n \geq 2$. If $x_1, x_2, \dots, x_n$ are non-negative real numbers, then

$$\frac{x_1 + x_2 + \dots + x_n}{n} \geq (x_1 x_2 \dots x_n)^{1/n}.$$

We saw the Principle of Proof by Induction in Theorem 0.3 and we also saw two variants in Corollary 0.5 and Theorem 0.6 (strong induction). I'd like to use this next proof to demonstrate another variant.

Proof. First consider the case where $n = 2^m$ for $m \in \mathbb{N}$. We'll prove this by strong induction on m . The case $m = 1$ is Theorem 4.15 which we've already proved (twice!).

Now suppose that the statement is true for n and consider the statement for $2n$. We will rewrite the arithmetic mean so that we can use our inductive hypothesis on the first half of the numbers (x_i with $1 \leq i \leq n$), and separately on the second half of the numbers (x_i with $(n + 1) \leq i \leq 2n$), writing

$$\frac{1}{2n} \sum_{i=1}^{2n} x_i = \frac{1}{2} \left(\frac{1}{n} \sum_{i=1}^n x_i + \frac{1}{n} \sum_{i=n+1}^{2n} x_i \right).$$

Then we have

$$\begin{aligned} \frac{1}{2} \left(\frac{1}{n} \sum_{i=1}^n x_i + \frac{1}{n} \sum_{i=n+1}^{2n} x_i \right) &\geq \frac{1}{2} \left((x_1 x_2 \dots x_n)^{1/n} + (x_{n+1} \dots x_{2n})^{1/n} \right) \quad \text{by the inductive hypothesis} \\ &\geq \sqrt{(x_1 x_2 \dots x_n)^{1/n} \times (x_{n+1} \dots x_{2n})^{1/n}} \quad \text{using the case } m = 1 \\ &= (x_1 x_2 \dots x_n x_{n+1} \dots x_{2n})^{1/(2n)}. \end{aligned}$$

This completes the proof for powers of 2. For all other numbers, we proceed by something resembling induction, but in the opposite direction.

Suppose the statement is true for n . We wish to prove the statement for $n - 1$.

Given $n - 1$ numbers with $n \geq 2$, let $x_n = \frac{1}{n-1}(x_1 + x_2 + \cdots + x_{n-1})$ i.e. the arithmetic mean of the others. Let's call this number M . Now since the statement is true for n , we have

$$\frac{x_1 + x_2 + \cdots + x_{n-1} + M}{n} \geq (x_1 x_2 \cdots x_{n-1} M)^{1/n}$$

The left-hand side simplifies to M and we have

$$M \geq (x_1 x_2 \cdots x_{n-1} M)^{1/n}.$$

Dividing both sides by $M^{1/n}$, we have

$$M^{(n-1)/n} \geq (x_1 x_2 \cdots x_{n-1})^{1/n},$$

and taking each side to the power of $n/(n - 1)$ gives

$$M \geq (x_1 x_2 \cdots x_{n-1})^{1/(n-1)}.$$

This completes the proof; the statement is true for $n - 1$.

Since every number is either a power of 2, or can be reached by starting from a power of 2 and repeatedly subtracting one, the statement is true for all $n \geq 2$. $\square$

Counterexamples

Providing a counterexample is the best method for *refuting*, or *disproving*, a conjecture. In seeking counterexamples, it is a good idea to keep the cases you consider simple, rather than searching randomly. It is often helpful to consider ‘extreme’ cases, in which something is zero, a set is empty, or a function is constant, for example. If you are relaxing one of the hypotheses of a theorem and contemplating whether the conclusion still holds, make sure to consider cases that contravene the relaxed hypothesis (else you already know that they won't provide the desired counterexample!)

For example, suppose it were claimed that the requirement for x and y to be non-negative could be removed from Theorem 4.15 if we simply put a modulus sign inside the square root:

Claim. *Let x, y be real numbers. Then $\frac{1}{2}(x + y) \geq \sqrt{|xy|}$, with equality if and only if $x = y$.*

Refutation. This claim is not true. A counterexample is $x = 1$, $y = -1$.

There is no need to expand with additional arguments about why the counterexample exists; providing a single counterexample is sufficient to disprove the claim.

4.4 Examples of problem-solving

In this section we discuss some example problems to illustrate aspects of problem-solving.

The following example explores how the images and preimages of set intersections behave. It is presented as a simple true/false question, to which it should be understood that we need to provide reasoning for our answer. So we first need to decide whether we think they are true or false (some experimenting and thought may be required); then if true, we should prove it,

and if false, we should provide a counterexample. A counterexample may well have been found anyway as part of our initial investigation to decide whether it is true or false, or conversely we may have gained insight into how a proof could work.

Example 4.17 (Images and pre-images)

Let $f: X \rightarrow Y$ be a function and let $A, B \subseteq X$ and $C, D \subseteq Y$. Are the following statements true or false?

- (i) $f(A \cap B) = f(A) \cap f(B)$,
- (ii) $f^{-1}(C \cap D) = f^{-1}(C) \cap f^{-1}(D)$.

Solution and commentary. For both of these it might be reasonable to consider some extreme cases in case we stumble across an immediate counterexample. In this situation, such an extreme case might be if A and B , or C and D , are disjoint. In that case the sets on the left hand sides are both the empty set (the image and the pre-image of the empty set are the empty set). In the first case this immediately suggests the possibility of a counterexample, since $f(A)$ and $f(B)$ could easily intersect. If the function were constant, for example (that is, if the function assigns the same output to every input), then $f(A) = f(B)$, so provided A and B are not themselves empty their intersection will not be the empty set.

So we claim (i) is false. A counterexample would be $f: \mathbb{R} \rightarrow \mathbb{R}$ defined by $f(x) = 3$ for all $x \in \mathbb{R}$. Then if $A = \{0\}$, $B = \{1\}$, then $f(A \cap B) = \emptyset$, but $f(A) \cap f(B) = \{3\}$.

We might reflect on whether anything can be salvaged (this is not part of the question, but is good practice). In particular, based on our experience finding the counterexample, we might suspect that $f(A \cap B) \subseteq f(A) \cap f(B)$. This is in fact true, and you might like to prove it.

For (ii), the same thinking does not yield a counterexample, since if C and D are disjoint then their pre-images are also disjoint. So we claim that (ii) is true and aim to prove it:

Proof. If $x \in \text{LHS}$, then $f(x) \in C \cap D$, so $f(x) \in C$ and $f(x) \in D$. Thus $x \in f^{-1}(C)$ and $x \in f^{-1}(D)$, and therefore $x \in \text{RHS}$. Conversely, suppose $x \in \text{RHS}$. Then $x \in f^{-1}(C)$ and $x \in f^{-1}(D)$, so $f(x) \in C$ and $f(x) \in D$, and therefore $x \in \text{LHS}$. So each side is a subset of the other, and the sets are therefore equal. $\square$

The next example relates to ideas that are covered in the Groups and Group Action course. This example is phrased in a way that is similar to a problem-sheet or exam-style question, with related sub-parts. The way the two questions in (ii) are phrased strongly suggests that there must be a difference between the case when n is prime and when n is arbitrary. A ‘hint’ is given at the end, and we should think about how we could relate that to the question that is asked.

Example 4.18 (Modular arithmetic)

Let $n \geq 2$ be an integer and let $\mathbb{Z}_n$ be the set of equivalence classes $\{\overline{0}, \overline{1}, \dots, \overline{n-1}\}$ defined by congruence modulo n on $\mathbb{Z}$.

- (i) Show that the operation $\otimes$ on $\mathbb{Z}_n$ defined by

$$\overline{x} \otimes \overline{y} = \overline{x \times y},$$

is well-defined, where $x \times y$ denotes standard multiplication on $\mathbb{Z}$.

- (ii) If $\bar{x} \neq \bar{0}$, a multiplicative inverse $\bar{y}$ has the property that $\bar{x} \otimes \bar{y} = \bar{1}$. Is there a multiplicative inverse for every $\bar{x} \neq \bar{0}$? What if n is prime?

[You may assume Bezout's lemma, which says that if integers a and b are coprime, there exist integers k and l such that $a \times k + b \times l = 1$.]

Solution and commentary. Firstly, it is helpful to recall the precise definition of the equivalence relation 'congruence modulo n '; that is $x \sim y \Leftrightarrow y - x$ is a multiple of n .

For (i), we should consider why the given definition might *not* be well-defined. The concern must be that the same equivalence class can be represented in terms of different x and y (e.g. $\bar{0} = \bar{n}$, etc.). Since the given definition depends on x and y themselves, it might give a different answer if we represent the elements $\bar{x}$ and $\bar{y}$ using different values of x and y . So, supposing $\bar{x}_1 = \bar{x}_2$ and $\bar{y}_1 = \bar{y}_2$, we need to show that $\bar{x}_1 \otimes \bar{y}_1 = \bar{x}_2 \otimes \bar{y}_2$. By definition of the equivalence classes $x_2 - x_1 = kn$ and $y_2 - y_1 = ln$ for some integers k and l . So

$$x_2 \times y_2 = (x_1 + kn) \times (y_1 + ln) = x_1 \times y_1 + (x_1 l + y_1 k + kln)n,$$

(we have omitted some of the $\times$ symbols to save space). Since the final term is a multiple of n this means that $\overline{x_2 \times y_2} = \overline{x_1 \times y_1}$, so indeed we have $\bar{x}_1 \otimes \bar{y}_1 = \bar{x}_2 \otimes \bar{y}_2$. Hence, $\otimes$ is well-defined.

For (ii), the question seems to suggest that things may be different depending on whether n is prime or not. So consider first a case where n is not prime and experiment a little to see if there are any simple counterexamples. If $n = 4$, there are only four equivalence classes, $\bar{0}$, $\bar{1}$, $\bar{2}$, and $\bar{3}$. We observe that $\bar{1}$ and $\bar{3}$ are their own inverse, but $\bar{2}$ does not have one. So we have found a counterexample: $n = 4$ and $\bar{x} = \bar{2}$. So the answer is No, there is not necessarily a multiplicative inverse for every $\bar{x} \neq \bar{0}$.

If n is prime, we might try to use the hint. If n is prime then for any $0 < x < n$, x and n will be coprime, so the lemma implies that there are integers k and l such that $x \times k + n \times l = 1$. But this means that $\overline{x \times k} = \bar{1}$ so, following the definition, $\bar{x} \otimes \bar{k} = \bar{1}$. So this $\bar{k}$, which we know exists according to the lemma, is the multiplicative inverse of $\bar{x}$. So if n is prime, the answer becomes Yes, every $\bar{x} \neq \bar{0}$ does have a multiplicative inverse.

This example has shown that if p is prime, then every non-zero element of $\mathbb{Z}_p$ has a multiplicative inverse. This goes part way to showing that $\mathbb{Z}_p$ is a finite *field* (meaning that it behaves in many respects similar to $\mathbb{R}$).

The final example below relates to things you will see in the Analysis II course. We are given a definition that involves a complicated-looking statement involving quantifiers, which is describing rigorously what it means for a function to tends to zero at infinity. We are being asked to apply this definition to two particular functions.

Example 4.19 (Limits)

A continuous function $f: \mathbb{R} \rightarrow \mathbb{R}$ tends to zero as $x \rightarrow \infty$ if

$$\forall \varepsilon > 0, \exists X \in \mathbb{R} \text{ s.t. } \forall x \in \mathbb{R}, \text{ if } x > X \text{ then } |f(x)| < \varepsilon.$$

Prove or disprove whether the following functions tend to zero as $x \rightarrow \infty$:

- (i) $f(x) = e^{-x}$;
- (ii) $f(x) = \cos x$.

Solution and commentary. Hopefully we have some immediate intuition (from the shape of their graphs, for example) that the first function does tend to zero and the second one doesn't.

For (i), we aim to show that the definition holds. Since the statement starts with $\forall \varepsilon$, we should start our proof by letting an arbitrary $\varepsilon > 0$ be given. Then we need to show that there exists an X such that for all $x > X$, $|f(x)| < \varepsilon$. Since the function $f(x) = e^{-x}$ is decreasing, and is always positive, this can be achieved by taking $X = -\ln \varepsilon$. Then for $x > X$, $|f(x)| = |e^{-x}| < |e^{-X}| = \varepsilon$. So we have shown that the definition holds, and $f(x) = e^{-x}$ does tend to zero as $x \rightarrow \infty$.

For (ii), we aim to show that the definition does not hold, so we need to prove its negation. Following the rules for how to negate quantifiers, that is,

$$\exists \varepsilon > 0 \text{ s.t. } \forall X \in \mathbb{R}, \exists x \in \mathbb{R} \text{ s.t. } x > X \text{ and } |f(x)| \geq \varepsilon.$$

(The original statement here is of the form ' $\forall \exists \forall P(x)$ ', where P is itself of the form $Q \Rightarrow R$, in which Q is ' $x > X$ ' and R is ' $|f(x)| < \varepsilon$ '. So the negated statement is of the form ' $\exists \forall \exists \text{ not } P(x)$ ', and $\text{not } P(x)$ has been expressed as ' Q and $\text{not } R$ '.) To see that this negated statement is true, we can observe that if $\varepsilon = \frac{1}{2}$ then for any $X \in \mathbb{R}$ there is a multiple of 2π , say $2\pi n$, that is larger than X , and for which $\cos 2\pi n = 1 \geq \varepsilon$. Hence $f(x) = \cos x$ does not tend to zero as $x \rightarrow \infty$.

4.5 General advice

When seeking to prove or disprove a result, the following suggestions may be helpful:

- Make sure you are clear about the hypotheses and conclusions.
- 'Unpack' any definitions and re-state what exactly it is you know and what it is that you need to show (either in your head or, if it's helpful, write it down).
- If you need to show something 'for all $\varepsilon > 0$ ', start with 'Let $\varepsilon > 0$ be given.'
- If you can't see a way to start, consider 'seeking a contradiction' and suppose the result is not true to give yourself more to work with.
- If you need to show uniqueness, suppose there are two of whatever it is, and try to show that they are equal.
- Look for extreme/simple cases as counterexamples.
- Don't be afraid to experiment, but have in mind what you're aiming for. If you're not making progress, try a different approach.
- Use sketches and diagrams to help gain intuition.

- If you get stuck, take a break. Look at it again with fresh eyes.
- Re-read your final proof. Be critical, and check that you are convinced by what you've written. (This is probably the most important of all these suggestions!)

A Notes for two in-person lectures

These are the notes for the two in-person lectures at the start of the year. The lectures won't follow this word-for-word, but will cover roughly this material in roughly this order.

For most courses, you can expect the in-person lectures to cover all of the material on the synopsis. This course is unusual, and the in-person lectures will cover some (but by no means all) of the definitions above, with some new connections between topics, and some more examples.

A.0 Lecture 0

Definition A.1

A natural number is a member of the sequence $0, 1, 2, 3, \dots$ obtained by starting from 0 and adding 1 successively.

We write $\mathbb{N}$ for the set of all natural numbers $\{0, 1, 2, 3, \dots\}$

Definition A.2

A set A is said to be a subset of a set S if every element of A is an element of S . We write $A \subseteq S$.

Definition A.3

The power set of a set A , denoted $\mathcal{P}(A)$ is the set of all subsets of A . It's a set of sets.

$\mathcal{P}(\mathbb{N}) = \{\emptyset, \{0, 1\}, \{n \in \mathbb{N} : n \text{ is prime}\}, \{3\}, \dots\}$ where the “...” is doing a lot of work.

A function from X to Y is an assignment of precisely one value in Y for each value in X . Can we have $f : \mathcal{P}(\mathbb{N}) \rightarrow \mathbb{N}$ by taking each subset to its “least element”?

Not quite, because this doesn't tell us what to do with $\emptyset$. The function f that takes subsets of the naturals to their least element and takes the empty set to 0 is well-defined, due to the Well-Ordering Principle, which states that every non-empty subset of $\mathbb{N}$ has a least element.

Definition A.4

Given a function f from X to Y , we say that the function is

- injective if whenever $f(x_1) = f(x_2)$ then $x_1 = x_2$.
- surjective if for every $y \in Y$ there is an $x \in X$ such that $f(x) = y$.
- bijective if it's both injective and surjective.

The function above is surjective but not injective.

Proof. for all n , $f(\{n\}) = n$ (so it's surjective) but also $f(\{n, n+1\}) = n$ (so it's not injective). $\square$

Is there a surjective function $f : \mathbb{N} \rightarrow \mathcal{P}(\mathbb{N})$?

No!

Proof. Suppose for contradiction that f is surjective, and consider the set

$$R = \{n \in \mathbb{N} : n \notin f(n)\}$$

f is surjective, so $\exists x \in \mathbb{N}$ such that $f(x) = R$. Is $x \in R$?

If $x \in R$, then by definition of R , $x \notin f(x)$. But $f(x) = R$ so $x \notin R$.

Similarly, if $x \notin R$ then $x \in f(x)$ so, x meets the condition to be in R , so x should be in R .

We have a contradiction. $\square$

So we cannot have a bijection from $\mathbb{N}$ to $\mathcal{P}(\mathbb{N})$.

Definition A.5

Two sets S and T are said to have the same cardinality if and only if there exists a bijection from S to T .

We've seen that $\mathbb{N}$ and $\mathcal{P}(\mathbb{N})$ do not have the same cardinality. Some cardinalities are assigned names or numbers;

- $\emptyset$ has cardinality 0.
- $\{1, 2, 3, \dots, n\}$ has cardinality n .
- $\mathbb{N}$ has cardinality $\aleph_0$.
- $\mathcal{P}(\mathbb{N})$ has cardinality $\mathfrak{c}$ (the cardinality of $\mathbb{R}$, proof in Analysis).

We can replace $\mathbb{N}$ in the proof above with any set S to show that S and $\mathcal{P}(S)$ do not have the same cardinality.

The sequence of sets $\mathbb{N}, \mathcal{P}(\mathbb{N}), \mathcal{P}(\mathcal{P}(\mathbb{N})), \dots$ is an infinite sequence of infinite sets, no two of which have the same cardinality.

A.1 Lecture 1

Definition A.6

Given sets A and B , the Cartesian product, denoted $A \times B$, is given by

$$A \times B = \{(a, b) : a \in A \text{ and } b \in B\}.$$

Definition A.7

Let X and Y be sets. A function f from X to Y is a subset of $X \times Y$ such that

$$\forall x \in X \quad \exists! y \in Y \quad \text{s.t.} \quad (x, y) \in f.$$

Definition A.8

A relation R on a set S is a subset of $S \times S$. If $(a, b) \in R$, we write aRb .

Definition A.9

A relation R on a set S is an equivalence relation if it is reflexive ($\forall x, xRx$), symmetric ($\forall x, y, xRy \Rightarrow yRx$), and transitive ($\forall x, y, z, (xRy \wedge yRz) \Rightarrow xRz$).
For an equivalence relation, we might write $a \sim b$ instead of aRb .

Example A.10

The equivalence relation on $\mathbb{R}$ given by $x \sim y$ if and only if $x - y \in \mathbb{Z}$.

(You might wonder which other sets $A \subseteq \mathbb{R}$ give an equivalence relation via $x \sim y \Leftrightarrow x - y \in A$.)

Definition A.11

Given an equivalence relation $\sim$ on a set S , and given $x \in S$, the equivalence class of x , denoted $[x]$, is the subset

$$[x] = \{y \in S : y \sim x\}.$$

Note that if $x \sim y$ then the sets $[x]$ and $[y]$ are equal (they have the same elements; prove this with the definition of an equivalence relation).

Definition A.12

Given an equivalence relation $\sim$ on a set S , the quotient set $S/\sim$ is the set of equivalence classes

$$S/\sim = \{[x] : x \in S\}$$

It's a set of sets.

Example A.13

Let $S = \mathbb{N} \times \mathbb{N}$ be the set of ordered pairs of natural numbers. Define an equivalence relation $(a, b) \sim (c, d)$ if and only if $a + d = b + c$.

Then the equivalence class $[(2, 3)]$ has elements like $(2, 3)$ and $(34, 35)$ and $(1728, 1729)$.

Informally the thing they have in common is the difference between the two components.

The quotient set $S/\sim$ has a lot in common with the set $\mathbb{Z}$ that you know and love.

We can define an addition operation on the equivalence classes by

$$[(a, b)] \oplus [(c, d)] = [(a + c, b + d)].$$

Note that this is well-defined because choosing a different representative of $[(a, b)]$ doesn't change the equivalence class on the right-hand side.

Example A.14

For the equivalence relation given by $x \sim y \Leftrightarrow x - y \in \mathbb{Z}$, the quotient set $\mathbb{R}/\sim$ is sometimes written as $\mathbb{R}/\mathbb{Z}$. The equivalence classes can each be represented as $[x]$ for some $x \in [0, 1)$.

The equivalence classes inherit an addition operation from $\mathbb{R}$, because $[x] \oplus [y] = [x + y]$ is well-defined.

But $[x] \otimes [y] = [x \times y]$ would not be well-defined. Counter-example; $[0.1] \otimes [0.1] = [0.01]$ but $[1.1] \otimes [0.1] = [0.11]$

Example A.15

For any function $f: X \rightarrow Y$, we can define an equivalence relation on X by $x \sim y$ if and only if $f(x) = f(y)$.

Reflexive; $\forall x, x \sim x$ is true because $\forall x, f(x) = f(x)$.

Symmetric; $\forall x, y$ if $x \sim y$ then $f(x) = f(y)$ so $f(y) = f(x)$, which means $y \sim x$.

Transitive; $\forall x, y, z$, if $x \sim y$ and $y \sim z$ then $f(x) = f(y)$ and $f(y) = f(z)$ so $f(x) = f(z)$ which means $x \sim z$.

For this equivalence relation, we can define a new function $F: X/\sim \rightarrow f(X)$ by setting $F([x]) = f(x)$.

This new function is well-defined because, no matter which representative of $[x]$ we choose, they all have the same value of f .

Claim: F is a bijection.

Proof. (Injective): if $F([x]) = F([y])$ then $f(x) = f(y)$ so $x \sim y$ and therefore $[x] = [y]$.

(Surjective): I changed the codomain from Y to $f(X)$. Every function is surjective onto its image. $\square$

As a specific example, let $X = \mathbb{R}$ and $Y = \mathbb{C}$ with $f(x) = e^{2\pi i x}$. Then the equivalence relation induced on $\mathbb{R}$ is precisely the one above where $x \sim y$ if and only if $x - y \in \mathbb{Z}$.

We have shown that $F: \mathbb{R}/\sim \rightarrow f(\mathbb{R})$ is a bijection, and here $f(\mathbb{R})$ is the unit circle and $\mathbb{R}/\sim$ is usually written $\mathbb{R}/\mathbb{Z}$.

Therefore $\mathbb{R}/\mathbb{Z}$ is in bijection with the unit circle.

As you progress through the degree, you'll see this strategy again and again, moving to a quotient set to simplify and understand many different objects. When a function isn't injective, we just quotient all our problems away!

B The Greek alphabet

A, α	alpha	H, η	eta	N, ν	nu	T, τ	tau
B, β	beta	Θ, θ	theta	Ξ, ξ	xi	Y, υ	upsilon
Γ, γ	gamma	I, ι	iota	O, o	omicron	Φ, ϕ or φ	phi
Δ, δ	delta	K, κ	kappa	Π, π	pi	X, χ	chi
E, ε or ϵ	epsilon	Λ, λ	lambda	P, ρ	rho	Ψ, ψ	psi
Z, ζ	zeta	M, μ	mu	Σ, σ	sigma	Ω, ω	omega

C List of Definitions

Definition 0.1 (Natural numbers)	3
Definition 0.2 (Less than or equal to)	3
Definition 0.7 (Divides)	6
Definition 0.8 (Prime number)	7
Definition 0.10 (Addition of natural numbers)	7
Definition 0.12 (Multiplication of natural numbers)	8
Definition 0.13 (Factorial)	8
Definition 0.15 (Binomial coefficient)	9
Definition 1.1 (Equal sets)	12
Definition 1.2 (Empty set)	12
Definition 1.3 (Subset)	12
Definition 1.4 (Proper subset)	12
Definition 1.7 (Bounded interval, unbounded interval, open interval, closed interval) .	14
Definition 1.8 (Power set)	15
Definition 1.10 (Cartesian product)	16
Definition 1.12 (Cartesian product of n sets)	16
Definition 1.13 (Union, intersection, complement, set difference)	17
Definition 1.14 (Disjoint sets)	17
Definition 1.21 (Finite set (recursive definition))	21
Definition 1.22 (Cardinality of a finite set (recursive definition))	21
Definition 2.1 (Relation)	23
Definition 2.5 (Reflexive, symmetric, anti-symmetric, transitive)	24
Definition 2.8 (Equivalence relation)	25
Definition 2.10 (Equivalence class)	25
Definition 2.11 (Quotient set)	26
Definition 2.12 (Partition)	26
Definition 3.1 (Function)	28
Definition 3.2 (Domain, codomain)	28
Definition 3.3 (Equal functions)	28
Definition 3.7 (Image, range, pre-image)	30
Definition 3.10 (Restriction of a function)	31
Definition 3.11 (Identity function)	31
Definition 3.12 (Injective, surjective, bijective, bijection)	31
Definition 3.15 (Finite set, cardinality of a set (with bijections))	32
Definition 3.16 (Function composition)	33
Definition 3.19 (Invertible function, inverse function)	35
Definition 3.24 (Left-invertible, right-invertible)	36
Definition 4.1 (And, Or, Not)	37
Definition 4.4 (Implies)	38
Definition 4.6 (Contrapositive, converse)	39